

<https://doi.org/10.1038/s41540-026-00657-8>

Calibrating tissue level PDE models of ligand dynamics using single cell and spatial transcriptomics data

Ali Daher¹✉, Dumitru Trucu² & Raluca Eftimie¹

Many physiological processes rely on cellular communication through ligand mediated signalling, where ligands secreted by cells diffuse through the extracellular space and bind to receptors to trigger downstream responses. Although reaction diffusion models capture these dynamics, parameter calibration often lags behind model development due to limited data that reflect the biological microenvironment and the lack of a unified framework for integrating such data into mechanistic models. We propose that single cell RNA sequencing and spatial transcriptomics provide a rich, abundant, and underused source of information for calibrating such models at tissue scale. We develop a computational pipeline that integrates these data to infer model parameters, combining finite volume solvers with bioinformatics preprocessing, Approximate Bayesian Computation (ABC), and gradient based optimization. Using two open source human skin datasets as case studies, we calibrate parameters governing the isoforms of Transforming Growth Factor Beta, a key regulator of tissue repair and fibrosis, and compare the resulting spatial concentration fields with local cell type distributions. Benchmarking on synthetic datasets helps evaluate the inference accuracy and shows the benefits of combining ABC with gradient based methods. Overall, the pipeline provides a flexible and rigorous approach for calibrating tissue level mechanistic models using modern transcriptomics data.

Cellular crosstalk plays a fundamental role in coordinating tightly regulated processes of tissue organization and repair. To make decisions, cells communicate through a variety of signaling modalities, depending on the physiological context and spatial range of communication. These include ligand-mediated paracrine and autocrine chemical signaling, contact-based juxtacrine signaling, neuronal synaptic communication, endocrine signaling in the form of hormones traveling through the bloodstream, and, as more recently unraveled, long-range mechanical traction forces through supracellular filaments that influence collective cell migration¹. Among these signaling pathways, paracrine and autocrine chemical signaling are particularly foundational within tissue microenvironments, mediating local interactions between neighboring cells (paracrine) or within the same cell (autocrine). In both of these modalities, a cell releases a ligand, such as a growth factor, which then diffuses through the extracellular space and binds to a complementary receptor, eliciting a downstream signaling cascade in the target cell. This form of short-range chemical signaling is central to numerous biological processes, including development, inflammation, wound healing, and tissue remodeling. Dysregulations in these pathways

have been implicated in a range of pathological conditions, including cancer, fibrosis, and chronic inflammatory diseases^{2–6}.

Given its central role in tissue regulation, chemically mediated cellular crosstalk has been the focus of extensive mathematical and computational modeling efforts^{7–9}. One of the most widely used approaches involves reaction-diffusion models, which characterize the production, diffusion, degradation, and receptor binding of signaling ligands using particle or continuum-based spatial or spatiotemporal partial differential equations (PDEs)^{10,11}. These equations can be used to compute the spatial ligand concentration fields across space and time. In addition, they can also be coupled to cellular dynamics equations capturing cellular proliferation, migration, and death, in order to incorporate the effects of chemical signaling on these processes.

For such models to be clinically or scientifically meaningful, their parameter values must be carefully calibrated in a physiologically relevant and data-informed manner. Historically, however, the parameter selection process in many of these mechanistic models has often lacked behind model development in terms of rigor, robustness, and biological fidelity¹².

¹Department of Mathematics, Marie and Louis Pasteur University, Besancon, France. ²Department of Mathematics, University of Dundee, Dundee, UK.

✉ e-mail: ali.daher@umlp.fr

Traditionally, parameter values have been calibrated using one or both of the following approaches. First, they are sometimes heuristically tuned through trial and error in computational simulations; for instance, by selecting values that avoid numerical instability or that produce model outputs consistent with broad biological expectations. Yet, the true strength of *in silico* models lies in their ability to probe biological mechanisms whose outputs are unknown or poorly understood; thus, reliance on expected behavior limits their exploratory power.

Second, in some cases, parameter values are indeed derived from *in vitro* experimental assays, such as ligand-receptor binding studies used to estimate binding kinetics; see refs. 10,13 for example. However, this approach presents its own limitations. *In vitro* conditions fail to fully recapitulate the complexity of the tissue's biological environment. Moreover, these experiments are often conducted under highly specific conditions, varying by species, cell type, or tissue context, making it difficult to generalize the resulting parameter values. In many cases, experimental data are only available for a subset of the required parameters, leaving the remainder uncertain. Even when data are available, reported values can span several orders of magnitude depending on the experimental design. For example, the binding rates for transforming growth factor Beta (TGF β) isoform TGF β 1 to the TGF β type II receptor (TGF β 1-R2) have been reported to range from 5.1×10^5 to 3.1×10^7 (Ms)⁻¹, depending on whether or not the receptors in experimental binding assay were immobilized¹³⁻¹⁷.

The gap between model development and parameter calibration raises important questions: what alternative sources of data can be leveraged to tune those parameter values, and what methodological framework can be employed to integrate such data into the model for robust parameter inference? In this paper, we demonstrate how high-throughput genetic sequencing data, specifically the combination of single-cell RNA sequencing (scRNA-seq) and Visium spatial transcriptomics (ST), can be leveraged to calibrate the parameters of reaction-diffusion-based models of chemical signaling between cells.

To this end, we develop a computational pipeline that integrates established tools with tailored implementations of optimization schemes in a novel sequential pipeline in order to bridge the gap between single-cell gene expression data and higher-scale model dynamics, thereby enabling robust parameter calibration. In Methods, we provide a detailed description of the sequential steps in the pipeline, along with a step-by-step case study in Results that demonstrates how high-throughput genetic data from a skin tissue sample undergoing normal wound healing and collected at day 30 post-incision¹⁸ can be leveraged to calibrate the parameters of the reaction-diffusion model during the wound remodeling phase. For the case study, we focus on the three human isoforms for transforming growth factor Beta (TGF β), a key regulatory cytokine that plays an essential role across all three stages of the skin wound healing (inflammation, proliferation, and remodeling) and homeostasis, as well as mediates a diverse series of cellular responses including growth suppression, changes in cell morphology, and expression of cell adhesion genes¹⁹, making it a crucial regulator of fibrosis, cancer progression, and immune modulation.

Finally, in the "Discussion," we discuss certain aspects of the computational pipeline, highlight its advantages as well as some of its limitations, and explore potential future work. Overall, our results demonstrate that high-throughput genetic sequencing data offer a rich, underutilized, and nowadays increasingly accessible (and often freely available) source of information in biology for parameter calibration. All the high-throughput datasets in this study were obtained from open-source databases, underscoring the broad accessibility of this approach. The computational pipeline developed draws on methods from bioinformatics, continuous optimization, numerical analysis, and Bayesian inference to facilitate a rigorous and robust parameter calibration process that is physiologically informed and data-driven. While showcased here in the context of skin (normal and abnormal) wound healing, the pipeline is broadly applicable to other biological systems characterized by complex, multicellular dynamics, such as cancer growth and tissue regeneration.

Results

Bioinformatics analysis

Our analysis uses the scRNA-seq and ST data from donor 3, collected at day 30 post-incision in the study by Liu et al.²⁰. The Uniform Manifold Approximation and Projection (UMAP) for Dimension Reduction embedding of the different clusters from the scRNA-seq analysis and their allocated cell type annotation are shown in Fig. 1a, and the cell count per type across the tissue sample is shown in the bar chart in Fig. 1b. The spatial distributions (q_{05}) for the selected cell types (fibroblasts, macrophages, and endothelial cells) within the tissue sample obtained from the combined scRNA-seq and ST analysis are shown in Fig. 2, and the standard deviations for the posterior distributions are shown in Supplementary Fig. 13. Using the CellPhoneDB toolkit to analyze the TGF β mediated cellular crosstalk from the transcriptomics data, the resulting ligand-receptor pair interaction strengths across the cell-type pairs and TGF β isoforms are shown in Fig. 3.

Application of the computational pipeline steps to the wound healing sample

Figure 4 shows the filtered scores of the reduced ensemble of parameter realizations (with $r \in [0.2, 0.65]$) when the rejection sampler algorithm is executed with $MC = 1 \times 10^5$ independent parallelized runs and an acceptance threshold of $\epsilon \geq 0.2$. Figure 5 demonstrates the evolution of the acceptance thresholds for the TGF β reaction-diffusion model throughout the Sequential Monte Carlo (SMC) Filtering populations, which reflects the gradual improvement in the correlation coefficients since the acceptance threshold for each population is computed from $\epsilon_t = q_{0.35}(\epsilon_{t-1})$. Figure 6 provides examples for how the posterior parameter distributions evolve and converge towards tighter distributions that are different from the starting one as the SMC algorithm progresses through the SMC populations.

Figure 7 demonstrates the composition of the two principal components (PCs) with the smallest eigenvalues in terms of the parameter realizations of the last SMC population when conducting principal component analysis (PCA). As described in *Parameter Calibration: The Computational Pipeline*, the parameters with strong contributions to these two PCs (coefficients in the corresponding eigenvectors) correspond to the stiff parameter linear combinations to which the output is most sensitive, and which must be constrained to achieve agreement between the model and data-derived interaction strengths. Furthermore, the results of the variance reduction analysis (not shown) corroborate the sensitivity analysis derived from PCA. The parameters identified from this PCA analysis as most sensitive were also parameters with relatively tight confidence intervals as identified from the likelihood-based analysis (see Supplementary Fig. 10).

After running the Approximate Bayesian Computation (ABC) rejection-sampler and SMC-ABC algorithms to prime the initialization seeds for the gradient-based stage, we performed a multi-start Newton-trust-region optimization to maximize the correlation (equivalently, minimize the negative log-likelihood (NLL)). The resulting maximum-likelihood estimate (MLE) was then used to compute the model-derived interaction strengths between cell types for each TGF β ligand-receptor pair using Eq. (16), which were in turn compared with those inferred directly from the transcriptomics data, as shown in Fig. 8. The two sets of interaction strengths exhibit a strong linear relationship with a Pearson's correlation coefficient of $r = 0.99$, indicating excellent agreement. These parameter realizations can also be used to numerically simulate the concentration fields of the ligands of interest (by solving for Eq. 15) in order to examine the patterns of their spatial distribution. Figure 9 shows the TGF β concentration fields obtained for each ligand using the same final parameter realization instance as in Fig. 8. The obtained concentration fields can also be compared with the spatial distribution of the different cell types to extract any colocalization patterns (see "Discussion").

Biological insights enabled by the computational framework

We first assess the transcriptional variation across the fibroblast subtypes present in the donor 3-day 30 wound-healing dataset. Figure 10a shows the distribution of fibroblast cells, annotated into the three subtypes, projected

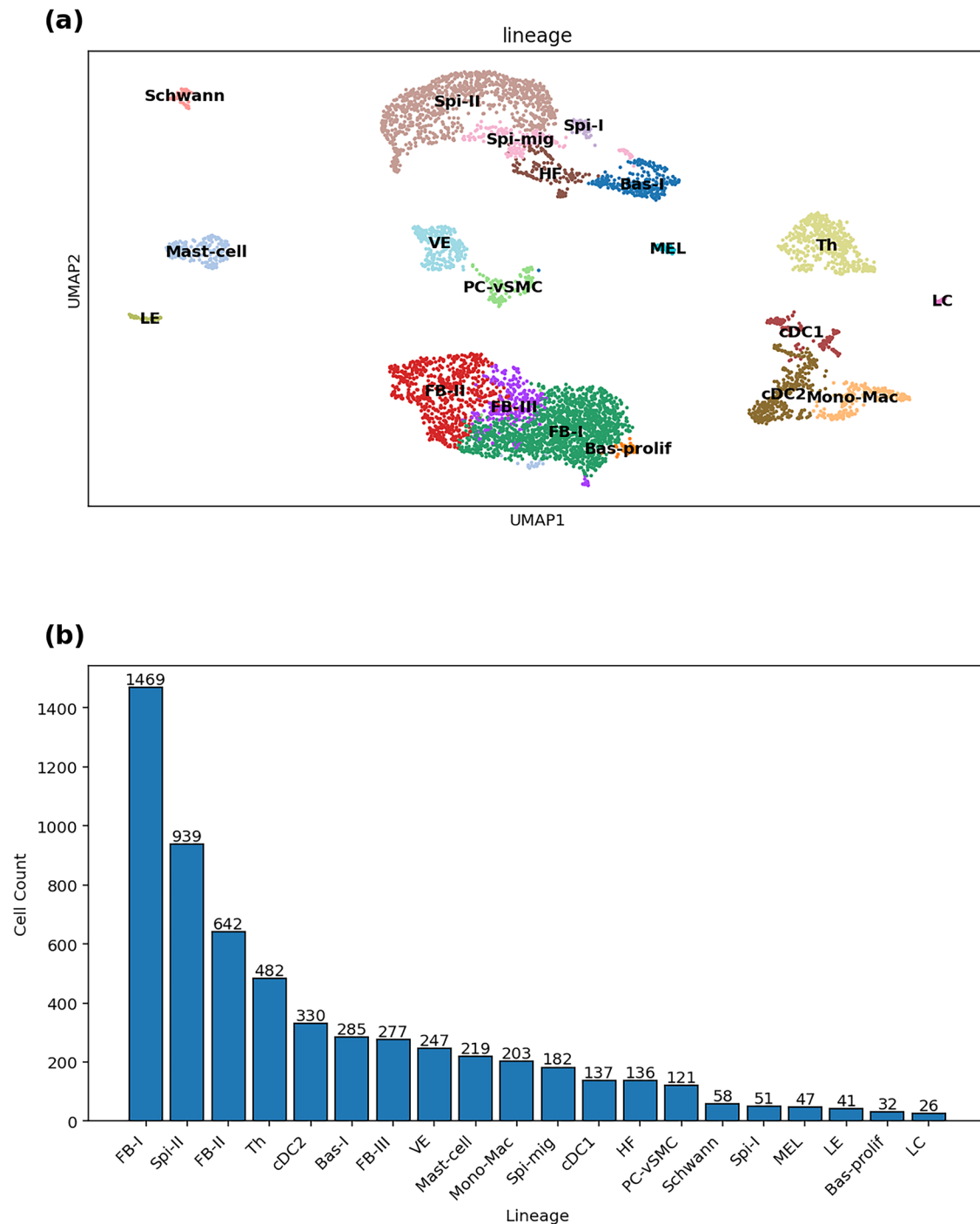


Fig. 1 | The processed raw scRNA-seq data (for donor 3 at day 30 from skin incision) in the study by Liu et al.¹⁸. a The UMAP embedding of the identified clusters with the corresponding cell type annotations. **b** The cell count per cell type. Bas/Spi basal/spinous keratinocytes, FB fibroblasts, HF hair follicle, LC Langerhans

cells, LE/VE lymphatic/vascular endothelial cell, MEL melanocyte, Mono-Mac monocytes and macrophages, PC-vSMC pericytes and vascular smooth muscle cells, Th T-helper cells, cDC conventional dendritic cells.

onto the first three PCs; although these components jointly capture less than 20% of the total transcriptional variance (Fig. 10b), the three fibroblast subtypes form distinguishable and separable clusters, which indicates that the subtypes are genetically and transcriptionally distinct, supporting the biological validity of treating them as functionally distinct populations in the model.

As evident in Fig. 11, FB subtype 2 exhibits a clear chemokine/pro-inflammatory transcriptional signature, along with notably lower extracellular matrix (ECM)-production scores than subtypes 1 and 3.

Furthermore, we hypothesized that FB subtype 1, which was annotated as mesenchymal in the original dataset, may contain a subset of myofibroblasts that were not distinctly labeled. The calculated myofibroblast scores support this interpretation: FB subtype 1 displays markedly higher myofibroblast-associated gene expression compared to the other subtypes. Based on these observations, we adopt a refined annotation in which the three subtypes are interpreted as mesenchymal/myofibroblast, pro-inflammatory, and papillary fibroblasts, respectively.

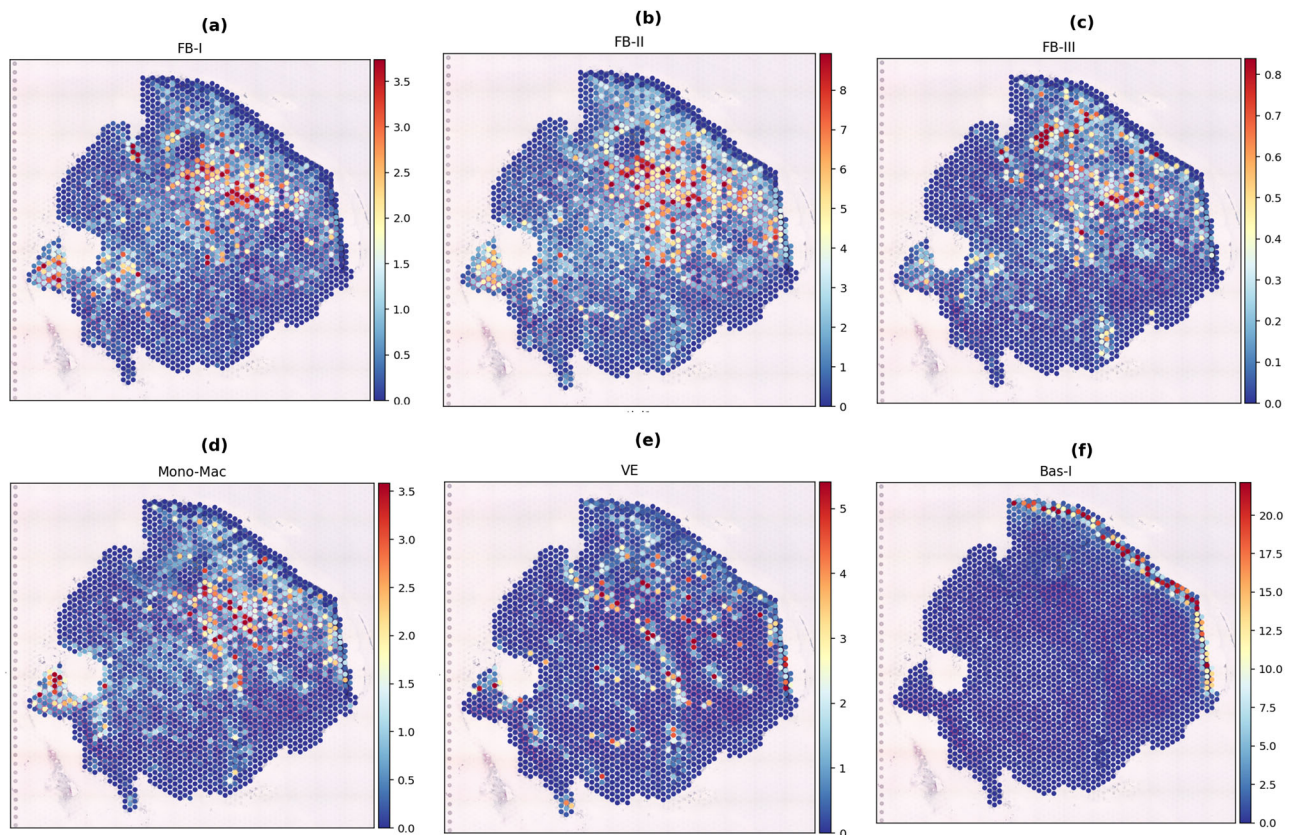


Fig. 2 | The spatial distribution across the Visium spots of six cell types in the normal-healing tissue sample (donor 3 at day 30 post-wound-incision)¹⁸. The cell counts represent the 5% quantile of the posterior distribution obtained from the

Cell2location analysis. **a–c** Fibroblasts types 1, 2 and 3 respectively, **d** Monocytes and macrophages, **e** Vascular endothelial, **f** Basal type 1.

Next, we compare the production rates of the different $TGF\beta$ isoforms across the different FB subtypes. Figure 12 shows the MLEs and corresponding confidence intervals, computed under the simplifying assumption of independent homeostatic noise, for the production rates of the three $TGF\beta$ isoforms across fibroblast subtypes. In addition, in Supplementary Note 6, we examine the effective amount of $TGF\beta$ receptors (moles/cell) between the normal day-30 wound-healing sample and the sample from acne keloidalis nuchae (AKN; subacute-chronic stage²¹).

Discussion

In this paper, we have developed a computational pipeline for the parameter calibration of cellular reaction-diffusion models from transcriptomics data. In particular, we looked at scRNA-seq and ST data as a potential rich, increasingly accessible, and perhaps underutilized source of data for the calibration of the reaction-diffusion models of ligands involved in chemical cellular communication. The computational pipeline developed draws on methods from bioinformatics, continuous optimization, numerical analysis, and Bayesian inference, integrating established computational toolkits with custom implementations and workflow adaptations in a unified framework to facilitate a physiologically informed and data-driven parameter calibration process, two important considerations if mathematical models are to have any substantial research and clinical utilities in the future.

The core contribution of this paper is the hierarchical combination of the three parameter calibration schemes into a single pipeline, as illustrated in Fig. 16, designed to overcome the individual limitations of each parameter calibration method. The computational pipeline developed in this study can be operated fully automatically, should the user choose. All code is freely available in an online repository. By relying solely on open-source data, the study demonstrates that the pipeline can be applied and calibrated using readily accessible information, making it

especially useful for research groups without direct access to high-quality experimental datasets.

This computational pipeline is primarily intended as a proof-of-concept demonstration that genetic data can be used to parametrize higher-scale model dynamics; hence, it was designed to be flexible, allowing modifications and refinements in its assumptions, the choice of parameters, mathematical formulations, and algorithms or toolkits. Examples of potential modifications include altering the mathematical formulations describing reaction kinetics (Eqs. 5–7) to accommodate more comprehensive binding dynamics, different formulations for the interaction strengths in Eq. (16) such that the production rate contributions could factor in ligands produced by neighboring control volumes (CVs), including more cell types in the analysis, and different choice of the toolkits or algorithms for extracting the interaction strengths. Below, we elaborate on these points and highlight several potential refinements for future work.

Our model relies on a quasi-steady-state approximation for $TGF\beta$ levels to simplify the PDE system, which is justified because during conditions such as the wound remodeling phase, $TGF\beta$ levels and cell population dynamics stabilize and $TGF\beta$ levels approach stable, non-zero concentrations, as described in Bioinformatics Preprocessing. In certain fibrotic tissues, $TGF\beta$ production can be elevated and unstable due to dysregulated feedback²², suggesting that the steady-state assumption may need to be revisited and the PDE formulation in the Mathematical Model adjusted accordingly. The pipeline could be extended to a dynamic (time-dependent) setting by introducing an appropriate time-marching scheme, such as the unconditionally stable Crank-Nicolson method, for the linear dynamic counterpart of Eq. (15), and computing the model-derived interaction strengths at multiple time points for comparison with the data-derived values. In this setting, the Jacobian of the growth-factor concentration fields with respect to the parameters, $\partial g \partial \theta$, as required for the gradient-based

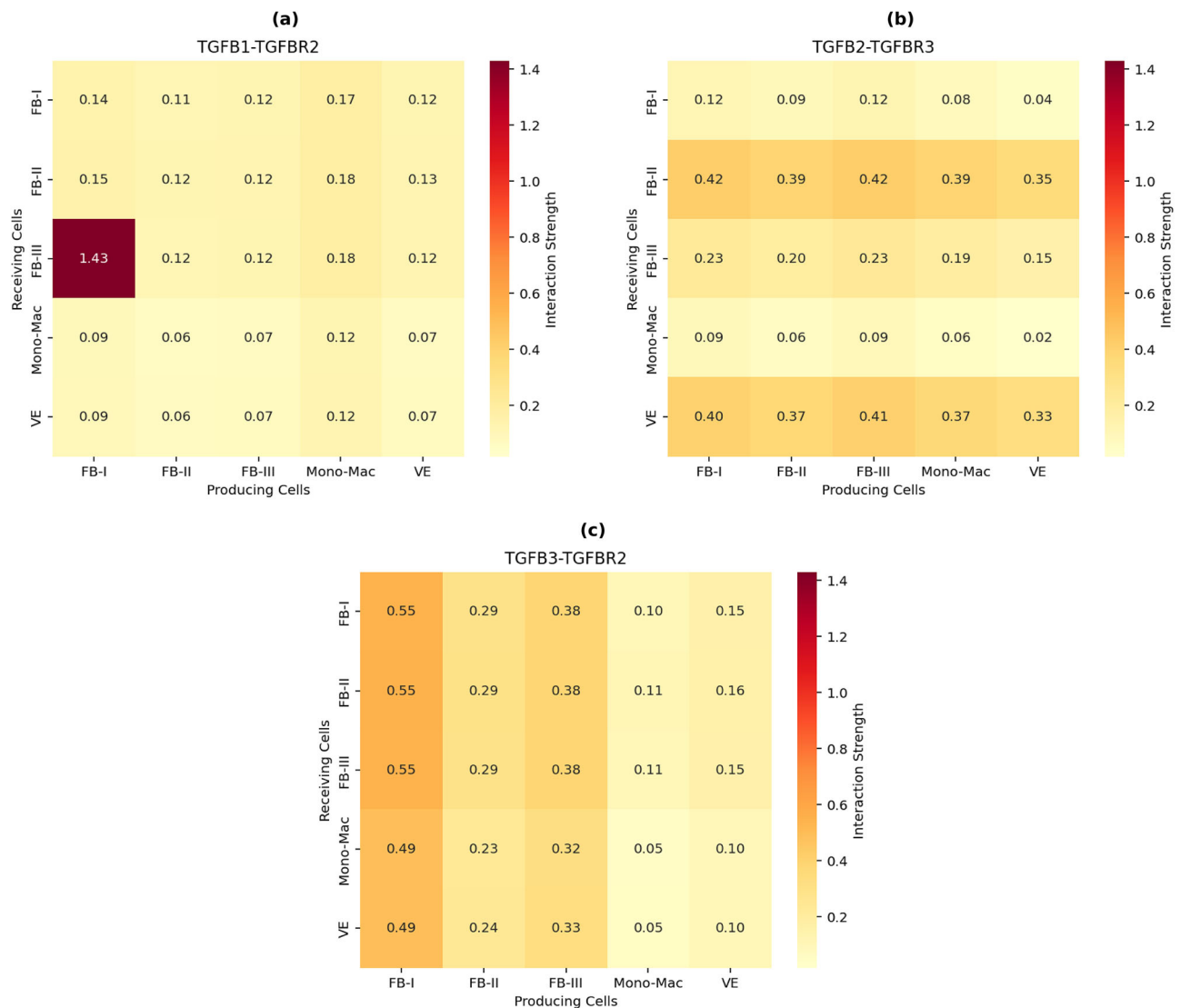


Fig. 3 | Interaction strengths between the cell types of interest (FB-I: Mesenchymal/Myofibroblast; FB-II: Proinflammatory; FB-III: Papillary; Mono-Mac: Monocytes and Macrophages; VE: Vascular Endothelial) and across the three TGFβ ligand isoform–receptor pairs, as inferred from the CellPhoneDB analysis

of the annotated scRNA-seq transcriptomic data for the wound-healing sample at day 30 post-incision. a TGFβ1–TGFβR2. b TGFβ2–TGFβR3. c TGFβ3–TGFβR2.

optimization stage, can be obtained, for example, using either forward or adjoint sensitivity analyses.

Another point of consideration is that the transcriptomics toolkits employed in the pipeline inherently introduce some uncertainty, even though such uncertainties are not routinely propagated in most bioinformatics pipelines. For clarity of exposition, the outputs of the Cell2location and CellPhoneDB toolkits were treated deterministically in *Bioinformatics Analysis*. We briefly outline how the principal sources of uncertainty arise, and how they may be quantified and propagated through our pipeline. Firstly, although tools for inferring interaction strengths like CellPhoneDB are deterministic, their output depends on cell-type annotations derived from the scRNA-seq clustering and cluster labeling. These annotations are influenced by different user-defined choices, including (i) the dimensionality reduction method (e.g., UMAP/t-SNE), which is stochastic, (ii) algorithmic hyperparameters (e.g., number of neighbors and resolution for UMAP), and (iii) the choice of clustering method (e.g., Leiden, Louvain), which can yield slightly different cell-type partitions; these, in turn, would result in different inferred ligand–receptor interaction strength matrices in Fig. 3, as well as different inferred gene signatures per cell type during the spot-gene deconvolution process. Because the mapping from the dimensionality reduction algorithm parameters to final interaction strengths is

complex and (for the most part) lacks a closed-form characterization, one principled approach would be to generate an ensemble of cell-type annotations by varying the algorithmic hyperparameters and then quantify the variability in the interaction strengths derived from CellPhoneDB (or toolkit of choice) across this ensemble. This would provide an empirical estimate of the uncertainty introduced at this stage.

Furthermore, in this study, we have relied on a point estimate of cell densities by using the 5% quantile (q_{05}) of the cell numbers per spot and cell type across the Visium hotspots obtained from deconvolution. The Bayesian inference package Cell2location, however, provides a full approximate posterior distribution (characterized by the mean, q_{05} , q_{95} , and standard deviation) generated via variational inference for the cell number of each type and spot to account for uncertainty in the cell type or number. In practice, the differences between q_{05} and the mean for the case studies in this paper were minimal, and the inferred distributions' standard deviations were small (see Supplementary Fig. 13), indicating tight distributions and low uncertainty. However, given the affine structure of the assembled finite volume (FV) system of equations as demonstrated in Eq. (14), these posterior ranges can be explicitly propagated through the reaction–diffusion solver using the delta method, under certain conditions and realistic assumptions (moderate standard deviations, smooth dependence on

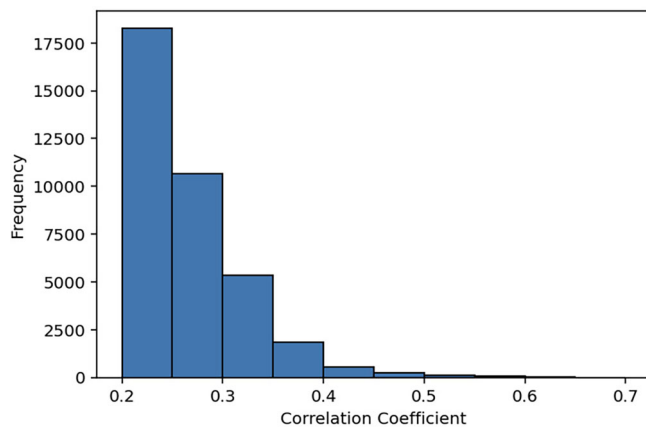


Fig. 4 | Acceptance scores (correlation coefficients) of the realizations obtained when applying the rejection sampler algorithm. Histogram showing the distribution of the correlation coefficients for the accepted runs when applying the rejection sampler algorithm ($MC = 1 \times 10^5$, $\epsilon = 0.2$) to the TGF β ligand reaction-diffusion parameter calibration.

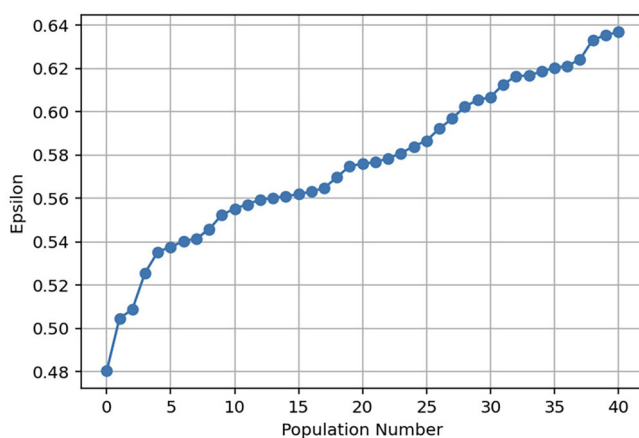


Fig. 5 | Evolution of the acceptance thresholds for the TGF β reaction-diffusion model throughout the SMC populations. Line plot depicting the acceptance threshold $\epsilon_t = q_{0.35}(\epsilon_{t-1})$ for each successive population.

parameters, symmetric posterior quantiles) from the data, as described in Supplementary Note 7.

An important implication of the type of such uncertainties is that noise across the inferred interaction strength data points is not independent. In scRNA-seq annotation, assigning a cell to one type necessarily reduces counts in another type; this introduces negative correlations in downstream interaction strengths. Similarly, in Cell2location deconvolution, the abundances of different cell types per spot are estimated from a shared gene-expression profile in that spot; allocating more counts to one type requires allocating fewer to others, which would in turn produce strongly correlated posterior uncertainties across cell types within the same spot. Therefore, the standard assumption of independent, homoscedastic, and additive noise, which we used to simplify the likelihood-based analysis and derive approximate confidence intervals, is unlikely to hold in practice. A fully rigorous likelihood-based treatment using Eq. (19) would entail modeling a correlated, heteroscedastic noise structure, which in turn require knowledge of cross-cell-type covariance structures (not automatically provided by current toolkits), or explicit modeling of uncertainty across all upstream steps, and would necessitate incorporating the dependency structure among residuals, for example, by using joint or conditional distributions rather than simple factorized likelihoods. Such an analysis is possible in principle but is substantially more involved and beyond the scope of the present study.

In such a case, the ABC scheme can still be advantageous by, for instance, sampling from the CellPhoneDB and Cell2location output ensembles without an explicit formulation of the noise structure and the corresponding likelihood formulation. Nevertheless, for the practical purpose of assessing parameter identifiability, we implement a likelihood-based objective under the simplifying assumption of homoscedastic, independent noise. This enables computation of confidence intervals around the MLEs, which in turn provides actionable insights into parameter identifiability and facilitates benchmarking of different calibration schemes, as demonstrated in *likelihood based analysis* and Supplementary Notes 3–5. Notably, ABC methods are particularly advantageous in such situations like ours, where the explicit computation of the true likelihood function is intractable: ABC, being a simulation-based approach, bypasses the need for explicit evaluation of the likelihood function. In practice, we have demonstrated that coupling ABC with a Newton-based method may yield MLE estimates with improved stability and, in many cases, tighter confidence intervals, reflecting enhanced local identifiability around the inferred solution. For parameters that exhibit practical parameter identifiability (small confidence intervals) in the synthetic datasets, the ABC-Newton method consistently yields tighter confidence intervals than the Newton-only approach, as demonstrated in Supplementary Note 5. In contrast, for parameters that exhibit no practical identifiability, both methods produce similarly wide confidence intervals, which is consistent with the inherent flatness of the likelihood surface in these directions. Furthermore, for both optimization schemes, the inferred MLEs lie close to the true parameter values across all parameters, and the true values fall within (or very near in a few cases) their corresponding confidence-interval bounds. Hence, the results of the benchmarking suggest that the ABC stages help steer the optimizer toward regions of higher curvature in the likelihood landscape. The Newton-trust-region solver can then potentially exploit this curvature through its Hessian-based updates, thereby improving local estimator precision. For non-identifiable or weakly identifiable parameters, confidence intervals remain broad under both initialization strategies, indicating that the ABC steps cannot resolve inherent structural or practical non-identifiability. While we focus here on a local Fisher-based identifiability analysis, the proposed pipeline is fully compatible with alternative approaches such as profile likelihood, which may be better suited for distinguishing structural versus practical non-identifiability or for capturing global, non-linear parameter dependencies²³.

We emphasize that these conclusions are derived from a specific set of experiments, namely, the synthetic interaction-strength dataset with known ground-truth parameters θ^* for the donor 3 wound-healing sample. We do not claim that ABC-initialized methods will universally outperform gradient-only approaches. In that regard, a broad, systematic evaluation of parameter inference strategies, similar to comparative studies such as^{24,25}, would be highly valuable. Nonetheless, the benchmarking analyses we have included in Supplementary Notes 4 and 5 represent meaningful progress in this direction and provide important context for the methodological choices in this work. In this paper, we have relied on a simplified kinetic representation of ligand-receptor binding and downstream signal propagation as a starting point for modeling TGF β signaling in our case studies. In reality, the binding kinetics are more complex, and in certain contexts, it may be important to incorporate these complexities into the model. For example, as explained in TGF β Binding Kinetics, TGF β binding involves the recruitment and formation of intermediate complexes before downstream signaling occurs. For simplicity, we have only modeled the interaction of the isoforms with their primary receptors. Similarly, our model assumes simple reaction dynamics and does not account for receptor saturation (either primary or recruited intermediates) or competition among isoforms. It is known that multiple ligand species can bind to multiple receptor species, resulting in competition. For example, of the 1735 (secreted) ligand-receptor pairs in the Fantom5 database³⁵, 72% of ligands and 60% of receptors bind to multiple species²⁶. This competition between multiple molecule species is prevalent and a crucial biophysical process, but it is generally ignored in cellular communication inference methods²⁷. In the particular context of TGF β modeling, both TGF β 1 and TGF β 3 bind to the

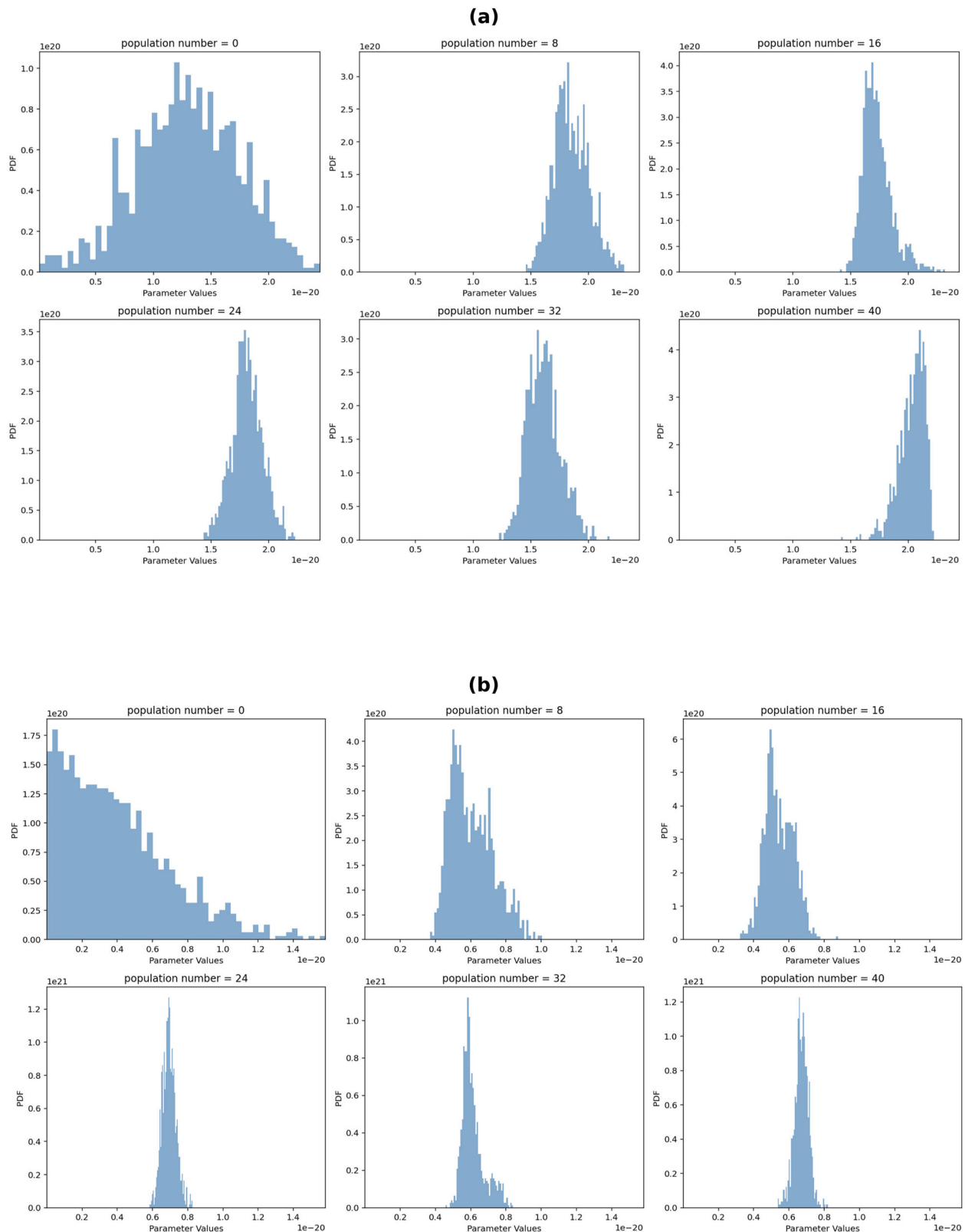


Fig. 6 | Tracing the evolution of the probability distributions for two of the TGF β model parameters across the different populations (0, 8, 16, 24, 32, and 40) of the SMC scheme. **a** R_1^{fib3} . **b** R_1^{fib2} .

same TGF β type II receptor, and all isoforms share the same intermediate protein for signal propagation. Furthermore, TGF β growth factors are secreted as latent complexes that require activation before binding to receptors²⁸. This activation step was not modeled explicitly, as during steady-state conditions, such as the remodeling phase, production and

activation rates equilibrate. However, when steady-state assumptions no longer hold, this mechanism would need to be explicitly included via additional equations. In the computational pipeline, cells are annotated by type based on the clusters identified in the scRNA-seq analysis, allowing differentiation between subtypes with meaningful genotypic and

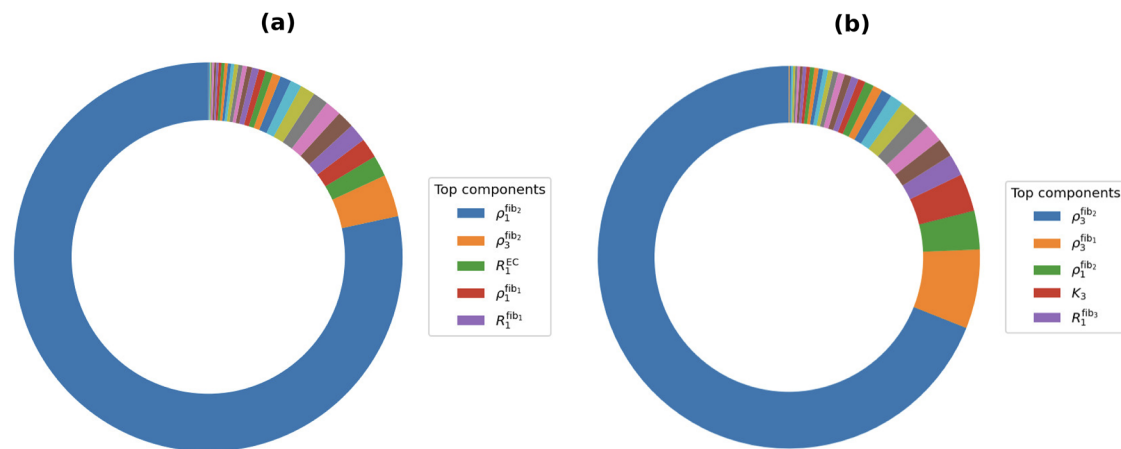


Fig. 7 | Constituent makeup of the two principal components (PCs) with the lowest eigenvalues (directions of least variance) from the covariance matrix of the final SMC population. **a** Makeup of the lowest-eigenvalue PC. **b** Makeup of the

second-lowest-eigenvalue PC. Parameters with prominent contributions indicate high constraint (output sensitivity) in the standardized space.

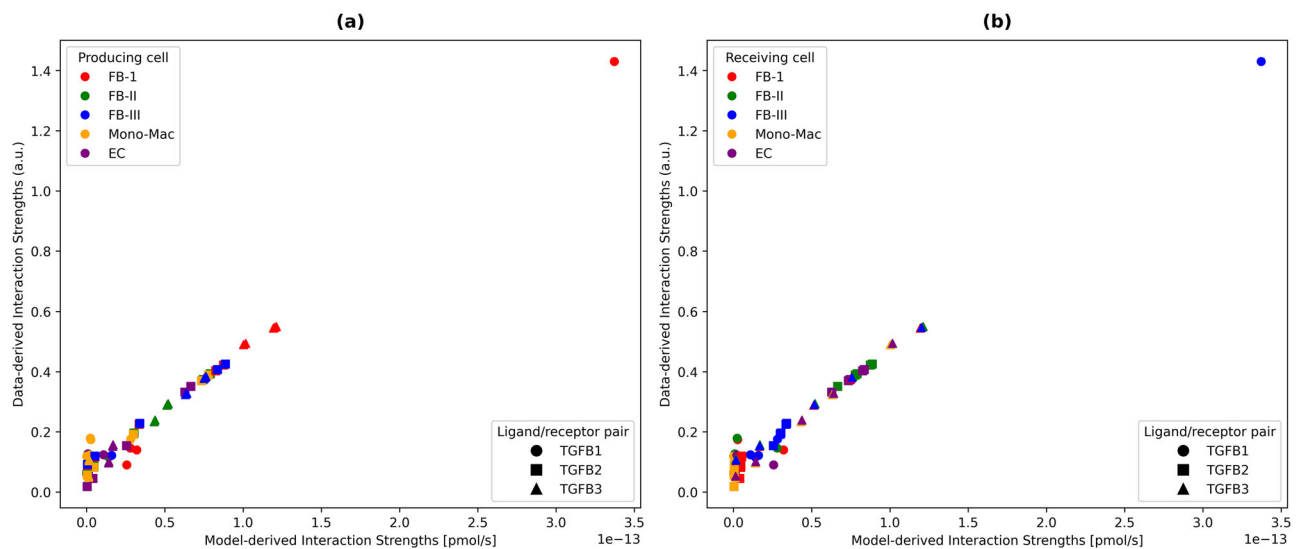


Fig. 8 | Model vs data-derived interaction strengths for the TGF β model for skin wound healing at day 30 post-incision, showing a strong linear relationship following parameter calibration ($r = 0.99$). The cells are color-encoded according to a TGF β secretion and **b** TGF β absorption. The genetic experimental interaction

strengths are derived directly from the transcriptomics data using bioinformatics tools as described in *Bioinformatics Preprocessing* (using the CellPhoneDB cell-cell communication inference tool) or Supplementary Note 1 (using COMMOT), while the model interaction strengths are calculated from Eq. (16).

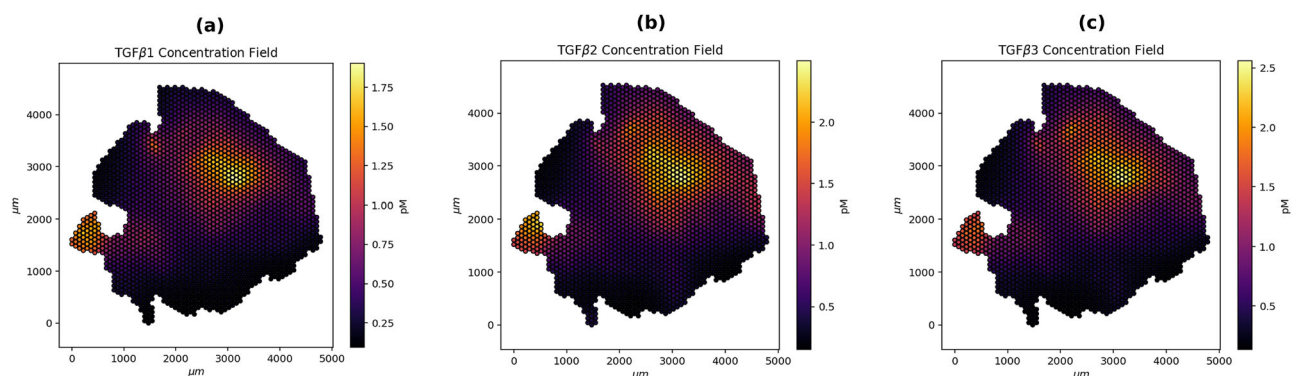


Fig. 9 | The numerically derived TGF β concentration fields for three ligands in the normal wound healing tissue sample at day 30 post-incision¹⁸ using the same parameter realization instance as in Fig. 8. The spatial distribution of the derived concentration fields for ligands: **a** TGF β 1, **b** TGF β 2, and **c** TGF β 3.

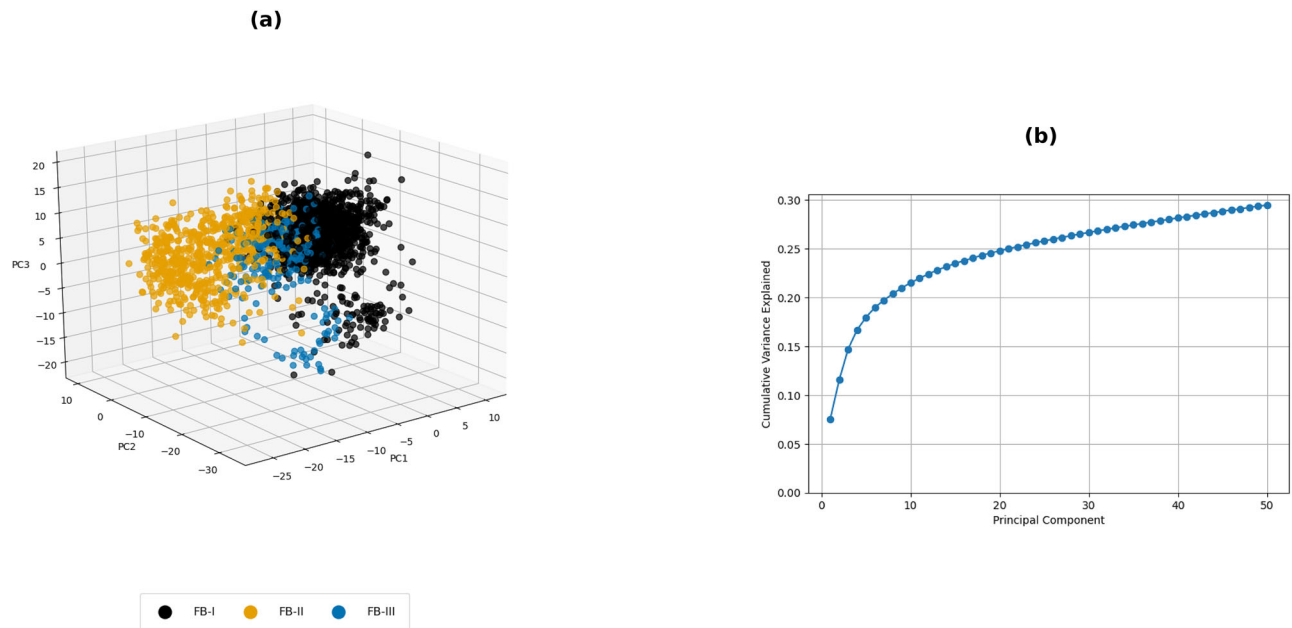


Fig. 10 | Principal Component Analysis (PCA) of scRNA sequencing data for fibroblast (FB) subtypes in the wound healing sample at day 30 post-incision. **a** Projection of the three annotated FB subtypes (FB-I: Mesenchymal/Myofibroblast;

FB-II: Proinflammatory; FB-III: Papillary) along the first three principal components. **b** Cumulative explained variance by the principal components.

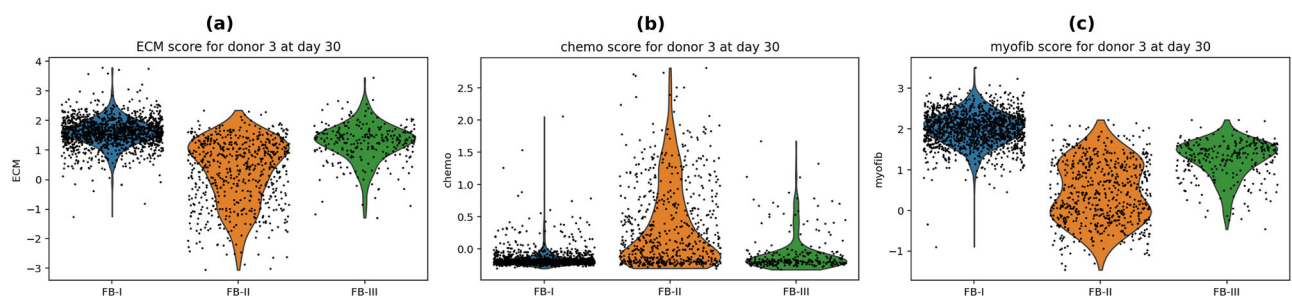


Fig. 11 | Gene biomarker scores across the three fibroblast (FB) subtypes (FB-I: Mesenchymal/Myofibroblast; FB-II: Proinflammatory; FB-III: Papillary) for the wound healing sample at day 30 post-incision. **a** ECM biomarkers. **b** Chemokine/inflammatory biomarkers. **c** Myofibroblast biomarkers.

phenotypic differences within the broader heterogeneous parent cell type (e.g., the three clusters of fibroblasts shown in Fig. 1a). However, within a given identified cluster or subtype, the current computational pipeline implicitly assumes that all cells share identical parameter values, such as ligand production rates or receptor numbers, which may not fully reflect biological variability. One direction for future work is to incorporate heterogeneity in production rates and receptor abundances across spatial locations, where, instead of assuming uniform parameters for all cells within a lineage, ST-derived (post-deconvolution) per-spot expression levels could be used to assign spot-specific parameters; this, in turn, would likely result in a more accurate characterization of the spatial growth factor concentrations fields, and spatial colocalization patterns with the different cell types, which are described below.

We note that the computational pipeline does not include rigorous parameter sensitivity²⁹ or structural identifiability analyses³⁰, as these are computationally demanding and difficult to incorporate into a fully automated workflow. Instead, our framework leverages an automated post-hoc analysis of the realization ensembles generated by the rejection sampler and SMC schemes, as well as relative confidence intervals, which in turn provide a robust and insightful assessment of parameter constraints, sensitivity, and identifiability. Importantly, they can be conducted post-hoc or in an

automated manner without requiring additional inference or variance-based simulations.

At this stage, the relationship between the interaction strengths derived from the model ligand concentration fields and those inferred from transcriptomics gene expression patterns has not been clearly established. Therefore, in this study, we adopt Pearson's correlation coefficient (r) as a simple, relevant, and interpretable measure of compatibility to be optimized, which would facilitate the likelihood-based confidence interval inference analysis described and implemented in Likelihood-based Analysis and Supplementary Note 3. In the future, as this relationship is better characterized, for example, through synthetic simulations, alternative evaluation metrics can be employed.

In this study, we inferred cell interaction strengths using CellPhoneDB, which estimates interactions based on the mean ligand and receptor expression between cell type pairs. These values were then compared to the values obtained from Eq. (16). We have used CellPhoneDB in this proof-of-concept study due to its ubiquity, extensive validation, widespread adoption, and relative ease of use^{6,21,31,32}. It should be noted, however, that CellPhoneDB does not incorporate spatial information regarding cell locations or distances between interacting cells, which is important since ligands and receptors often interact within limited maximal diffusion spatial ranges^{27,33}. Nonetheless, even without explicitly factoring in spatial information, it has

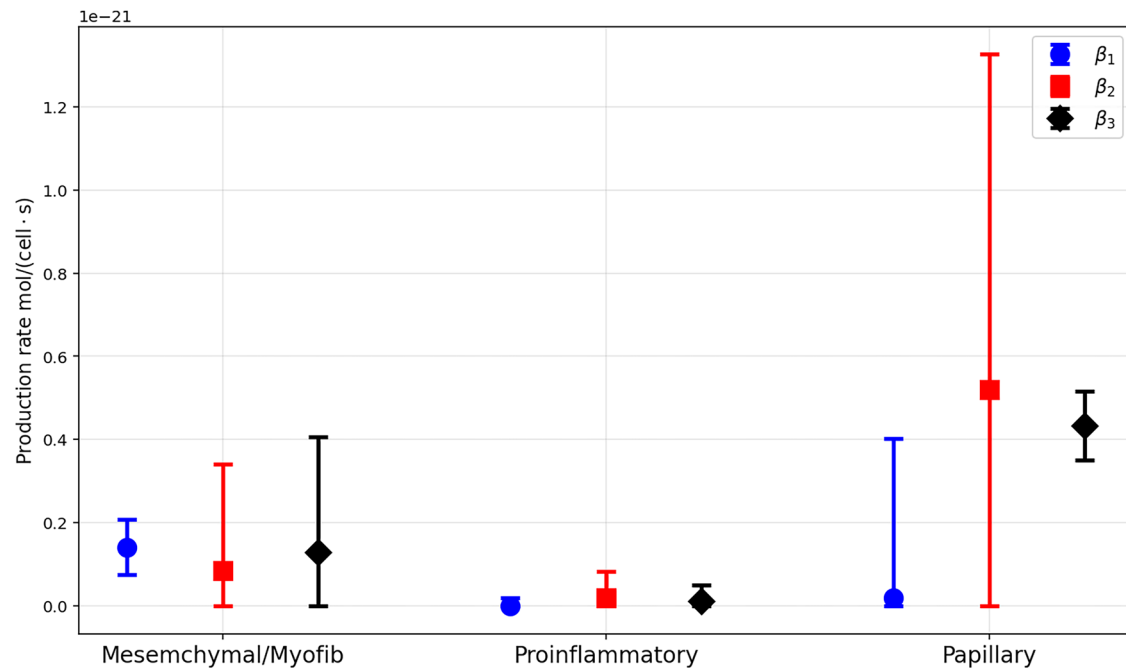
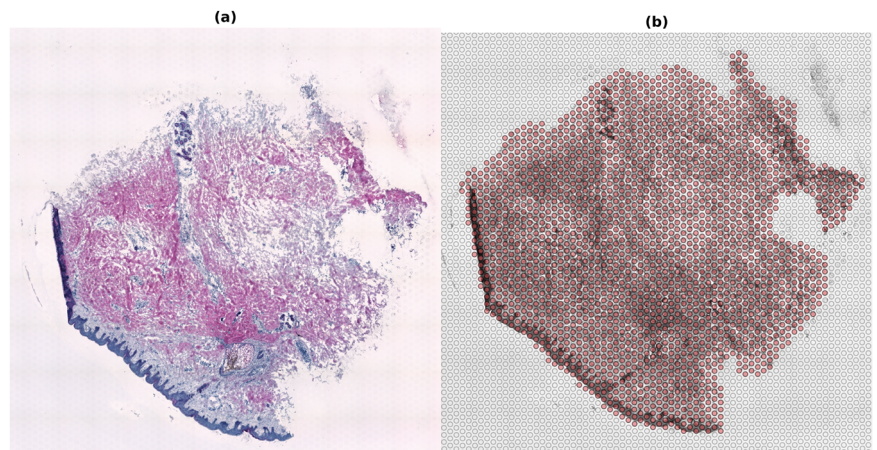


Fig. 12 | Comparison of inferred TGF β production rates (ρ_{fb}) across fibroblast (FB) subtypes for the wound healing sample at day 30 post-incision. Maximum Likelihood Estimators (MLEs) with 95% confidence intervals (error bars) from Fisher information, for the three TGF β isoforms (β_1 , β_2 , β_3) across FB subtypes.

Fig. 13 | Tissue sample representing normal wound healing for donor 3 at day 30 post-incision (from GEO accession GSE241132; originally published by Liu et al.¹⁸). **a** Histological tissue sample. **b** Corresponding Visium spots (red, hexagonal grid) overlaid on the tissue sample.



been demonstrated that CellPhoneDB (along with CellChat and a few other packages) can achieve overall very good performance in terms of consistency with spatial tendency and software scalability when compared to other available tools for dissecting cell-cell communications^{34,35}. Notably, a few other recent toolkits and extensions to existing ones have been very recently developed to incorporate spatial information when inferring cellular communication and interaction strengths from ST data^{27,33,36}. To demonstrate how our computational algorithm can accommodate the employment of other bioinformatics tools to derive the interaction strengths, we employ COMMunication analysis by Optimal Transport (COMMOT), which accounts for the competition between different ligand and receptor species as well as spatial distances between cells²⁷, to infer the cell-cell type interaction strengths from a sample of acne keloidalis²¹ in Supplementary Note 1. While existing bioinformatics tools for inferring cell-cell communication from transcriptomic data share similar underlying frameworks, each emphasizes different aspects of interaction networks, and their complementary strengths make them valuable when used together; comparative evaluations of these tools have been conducted elsewhere^{34,35}

and are beyond the scope of this study. Furthermore, in *Biological Insights Enabled by the Computational Framework* and Supplementary Note 6, we demonstrate how a post-hoc analysis of the parameter values obtained from the calibration pipeline could potentially provide valuable insights into the biophysical mechanisms of normal physiological processes, as well as those underlying disease progression, respectively. In Fig. 12, despite the relatively wide intervals, a clear pattern emerges: pro-inflammatory fibroblasts exhibit the lowest TGF β production across all isoforms. This is consistent with the biological timeline, as the inflammatory phase peaks early in days 2–5 after wound incision and largely subsides by the remodeling phase at day 30, at which the histological sample was taken. In addition, it is generally thought that mesenchymal fibroblasts dominate the remodeling phase, with papillary fibroblasts playing a supportive role, an interpretation supported by the larger number of mesenchymal versus papillary fibroblasts observed in our analysis (Figs. 1b and 10a). Nevertheless, Fig. 12 indicates a surprisingly strong activation of papillary fibroblasts, suggesting a potentially underappreciated contribution to the remodeling of the wound despite their small numbers. Of particular note is the relatively high TGF β 3 production rate

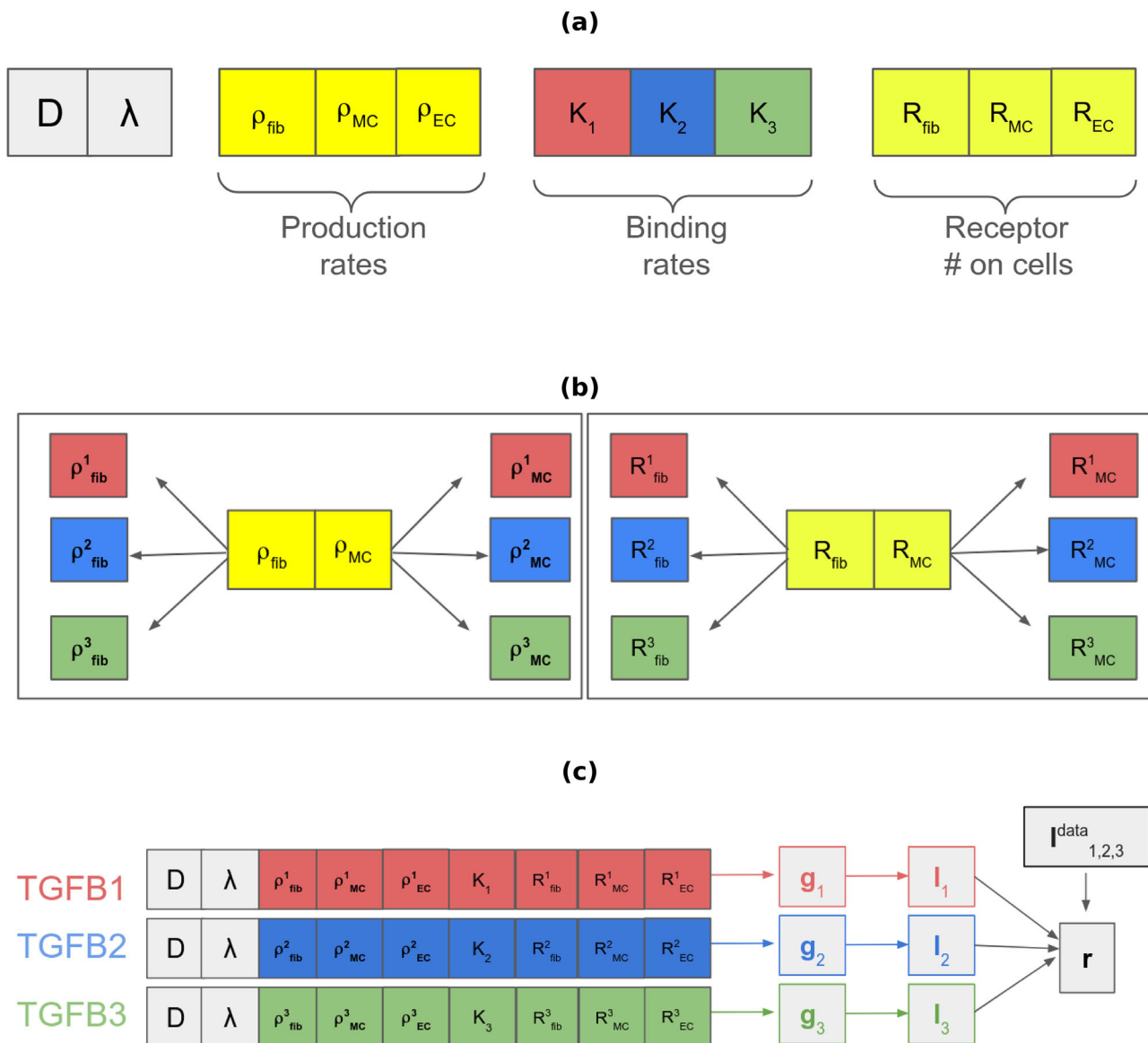


Fig. 14 | The parameter assembly framework for a particular parameter realization when modeling the reaction-diffusion dynamics of the TGFβ isoforms across the five cell types (the three fibroblast subtypes are represented here by one fibroblast type for brevity). a A multi-dimensional parameter realization is sampled from the prior distribution. **b** The production rates and number of receptors per cell are randomly partitioned across the three isoforms; the figure only shows the partitioning for the fibroblasts and macrophages, with similar portioning conducted for

the EC cell type. **c** The parameters are reassembled such that each isoform gets its own set of parameters. The D and λ_g parameters are assumed to be common across all isoforms. For each isoform, the FV linear system is solved to obtain the ligand concentration field, from which the model interaction strength matrix is obtained. The interaction strengths across the cell types and isoforms are then compared to the bioinformatics-derived interaction strengths data to obtain Pearson's correlation coefficient as a measure of agreement between the two datasets.

inferred for papillary fibroblasts (along with the tight confidence intervals): TGFβ3 is generally recognized as an anti-fibrotic growth factor that modulates ECM deposition and scar formation³⁷. This finding may therefore point to an important regulatory role of papillary fibroblasts in balancing ECM deposition and breakdown during the remodeling phase in order to limit fibrosis and promote more regenerative outcomes. As another example of the capacity to derive biological insight from the parameter inference procedure, we compare the inferred number of receptors in the sample of a normally healed wound vs acne keloidalis (Supplementary Note 6), and corroborate the hypothesis that fibrotic scarring is promoted as a result of increased receptive signaling and not only increased ligand production. The numerically obtained concentration fields can be compared with the spatial distribution of different cell types to identify potential colocalization patterns. For example, in Fig. 9, two clusters of high TGFβ concentrations are visible, as indicated by the yellow regions in the heatmap. Similarly, Fig. 2

shows two regions with high fibroblast and macrophage densities that are colocalized. The two clusters of high TGFβ concentrations coincide with the two clusters where fibroblast and macrophage cells are colocalized, potentially indicating increased TGFβ activity and thus cellular crosstalk between these cells.

It merits noting that analyzing cell-cell communication networks and interaction strengths from transcriptomics data to derive transcriptomics-based interaction strengths, as is routinely done using tools such as CellPhoneDB, CellChat, and COMMOT, relies on mRNA expression levels of genes encoding ligand and receptor proteins to identify likely ligand-receptor pair interactions. However, actual ligand-receptor physical interactions occur via proteins released and recognized by cells. Accordingly, mRNA expression levels of ligands and receptors serve as proxies for the actual physical interactions that occur at the protein level. However, mRNA levels may not necessarily directly correlate with the active protein levels due

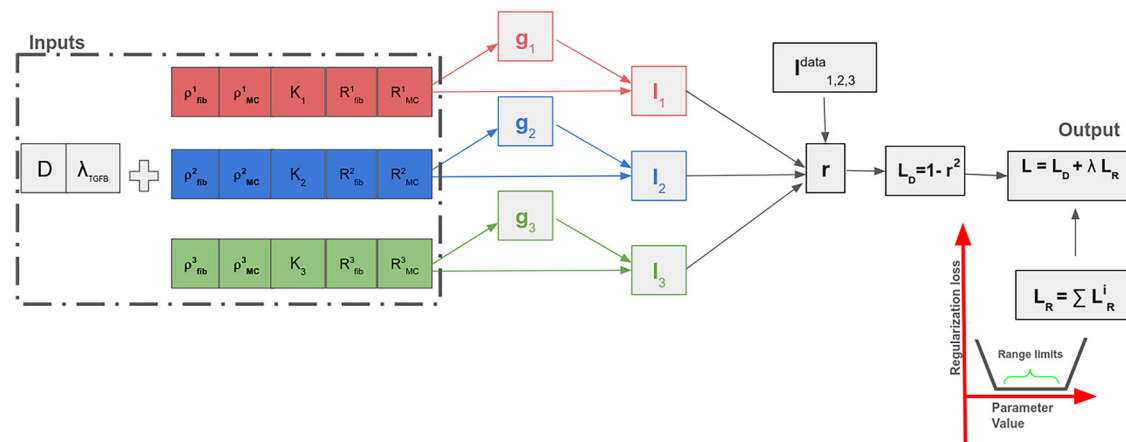


Fig. 15 | A simplified schematic demonstrating how the input parameter realizations are mapped through a series of differentiable mathematical operators to obtain the total loss function, which is comprised of the weighted sum of the

fitting and regularization loss functions. Using gradient-based algorithms, the effects of the parameters on the loss function can be traced back to update the parameters and minimize the loss function.

to myriad factors, including post-transcriptional and post-translational regulations and modifications, receptor protein trafficking and recycling, ligand protein degradation (decay), ligand protein activation from the latent forms in which proteins are released (for example via collagenase cleavage), and interactions with the ECM that modulate ligand availability; the last two being particularly relevant for the TGF β protein analyzed in this study.

Nonetheless, many prominent cell signaling studies have successfully relied on similar transcriptomic inferences to elucidate cell-cell communication patterns and interaction strengths (see³² for example). Furthermore, as previously discussed, the most prominently used tools for dissecting cell-cell interactions from transcriptomics data (CellPhoneDB, CellChat, and COMMOT) have produced consistent results in terms of relative interaction strengths between ligand-receptor and cell-cell pair combinations, despite employing different algorithms and frameworks. This consistency suggests that, while imperfect, these derived interaction strengths capture biologically meaningful patterns. These transcriptomics-based tools have been used here because they are ubiquitously employed in bioinformatics studies, and they provide the most comprehensive and accessible mapping of ligand-receptor interactions available, in addition to transcriptomics data being abundantly available, such as in public repositories. On the other hand, spatial quantification of proteins at a comparable scale remains technically infeasible with current technologies. Given that our mathematical model accounts for ligand-receptor interactions at the proteomic level through physical parameters such as diffusion coefficients, degradation rates, and binding kinetics, proteomic-derived interaction strengths (e.g., from multiplex protein imaging or spatial proteomics) could be directly incorporated into the pipeline in the future as these technologies continue to advance and become more accessible.

Methods

The first stage of the pipeline presented in this study, as detailed in “Bioinformatics preprocessing,” employs conventional bioinformatics methods and existing tools to extract two key data components from the sequencing scRNA and ST datasets. First, it computes the spatial cell density distributions of the relevant cell types across the tissue domain. Second, it estimates the ligand-mediated interaction strengths between different cell types and ligand-receptor pairs of interest, which would then serve as the experimentally derived reference data against which the model outputs can be quantitatively calibrated.

In the “Mathematical Model,” we formulate an FV scheme to compute the ligand concentration fields. The spatial discretization (mesh) of the FV solver mirrors the hexagonal grid used in the Visium ST experimental platform, which enables direct comparison and data integration between the

model outputs and the experimental data. For each multi-dimensional realization in the parameter space, the FV scheme is used to solve for the spatial concentration field of the growth factor(s) of interest. The input parameter realization, along with the numerically computed growth factor concentration fields, can be used to calculate the interaction strengths across all the relevant cell type pairs and ligand-receptor combinations of interest; these, in turn, can be directly compared to the experimentally derived interaction strengths in “Bioinformatics preprocessing.”

Next, in “Parameter initialization and assembly,” we conduct a thorough review of the existing literature to extract the relevant parameter values of the reaction-diffusion model by examining both experimental and computational studies. These values are used as a starting point to construct an informed prior distribution of the model parameters, from which an initial ensemble of realizations can be generated and subsequently updated following the calibration process.

All that remains is a framework through which the model-derived and the transcriptomics data-derived interaction strengths can be directly compared, with their differences systematically minimized and mapped back to the underlying model parameters. In “Parameter calibration schemes,” we discuss three relevant parameter calibration strategies: the rejection sampler scheme, the SMC scheme, and gradient-based optimization schemes. The first two fall under a class of computational Bayesian inference techniques known as ABC, whereas the third category encompasses standard continuous-optimization algorithms such as batch gradient descent, ADAM, and Newton-trust-region methods. We outline the advantages and limitations of each method in the context of calibrating parameters in the reaction-diffusion model, and in “Parameter calibration: the computational pipeline” we propose an integrated calibration pipeline where the three methods are applied sequentially in a hierarchical framework, such that the output of one method as the input for the next, in order to harness their individual strengths and mitigate their respective drawbacks.

Bioinformatics preprocessing

In this part, we describe how high-throughput genetic sequencing data, namely, the combination of scRNA-seq and ST datasets, can be effectively processed to extract relevant information that can be used in the subsequent numerical analysis and parameter calibration sections. Specifically, we outline how established bioinformatics methodologies and existing tools can be employed to compute two key quantities: (1) the spatial distribution of cells across the cell types of interest, which is needed to set up the FV numerical scheme, and (2) the interaction strengths for each ligand-receptor and cell-cell pair, approximated from the mRNA co-expression levels of the ligand in the secreting (sending) cells and the corresponding receptor in the

receiving cells; these interaction strengths then serve as the reference data against which the model output is compared and calibrated.

In a study on human skin wound healing, Lie and colleagues¹⁸ have induced wounds in healthy donors and collected wound-edge tissues at the three stages: inflammation (day 1), proliferation (day 7), and remodeling (day 30), monitoring the molecular and cellular dynamics of human and skin wound healing at each stage at an unprecedented resolution by performing both scRNA-seq and ST (Visium) experiments on the samples from two donors at each of the wound-healing stages¹⁸. In this study, we focus on analyzing the sample datasets for one of the two donors (donor 3) at day 30 post-wound incision, when the wound has already entered the remodeling phase. We choose to analyze the tissue sample at day 30 because experimental and computational studies^{10,38,39} suggest that during this period, the TGF β levels and cell population dynamics stabilize and TGF β levels approach stable concentrations. This means that we can simplify the governing equations and assume quasi-steady state dynamics for the TGF β field, governed by a time-invariant reaction-diffusion system ($\partial g/\partial t \approx 0$, where g is the TGF β concentration), while retaining the flexibility to extend the framework to fully time-dependent dynamics as needed, as outlined in the “Discussion.”

The goal of the scRNA-seq data processing step is to derive an annotated list, organized by cell type, of gene expression profiles (expressed in gene expression units) for the cells in the tissue sample, starting from the raw, unannotated gene expression matrix. The processing workflow follows largely standardized protocols and aligns with those employed in previous scRNA-seq studies aimed at characterizing the cellular and molecular landscapes of skin tissues and lesions^{18,21,40}:

1. **Filtering:** Low-quality cells expressing <400 or >8000 genes, $\geq 20\%$ mitochondrial genes, and <500 gene counts, as well as genes expressed in fewer than 10 cells, were excluded. Potential doublets detected using Scrublet⁴¹ were also excluded.
2. **Data Normalization and Scaling** were performed using the Scanpy package⁴² in Python.
3. **Dimensionality Reduction and Clustering:** The top 4000 variable genes were used for PCA using the Seurat flavor from Scanpy⁴², and the top 50 components (with the largest eigenvalues) were subsequently used in the clustering analysis. UMAP and k-nearest neighbor graph were generated using the Scanpy package. The major clusters were identified using the Leiden igraph-based algorithm in Scanpy (resolution = 0.8, neighbors = 15, and $n_pchs = 50$).
4. **Mapping Clusters to Cell Types:** Differentially expressed genes within each cluster were determined using the Wilcoxon method in Scanpy. To label each cluster by the cell type, the 50 most highly expressed genes in each cluster were determined and mapped against the differential gene expression (DGE) annotation file provided by the original study¹⁸; these, in turn, were originally obtained by using the “FindTransferAnchor” and “MapQuery” functions in Seurat^{18,43}.

Unfortunately, scRNA-seq data lack spatial context and do not provide information regarding the physical locations or distributions of individual cells within the tissue. To overcome this limitation, many high-throughput transcriptomic studies now integrate scRNA-seq with ST, most commonly using the Visium platform⁴⁴, which enables spatially resolved gene expression profiling. In the most widely used commercial implementation of Visium, a tissue section is mounted on a glass slide patterned with an array of 55 μm -diameter spots arranged in a hexagonal grid. Each spot captures mRNA from nearby cells, allowing the construction of a spots-by-genes expression matrix that retains spatial information. Figure 13a shows a histological image of normally healing tissue for donor 30 at day 30 post-incision, while Fig. 13b displays the corresponding Visium spots arranged in a hexagonal grid and overlaid on the tissue section, illustrating how spatial gene expression is measured across the wound bed¹⁸.

To map the spatial positions of the identified cells from the scRNA-seq analysis in their native environments (i.e., the histological tissue sample in Fig. 13), we use a computational toolkit called Cell2location⁴⁵. Cell2location

employs a Bayesian hierarchical model that decomposes the ST data into contributions from each of the identified cell types, estimating the proportions (or, in case prior information about the cell numbers is provided, the absolute numbers) of cells for each cell type at each Visium spot. In brief, Cell2location uses the top 100 marker genes for each cell type, obtained from the annotated (by cell type) cell-gene count matrix constructed from the scRNA-seq analysis stage, to train a negative binomial regression model and estimate reference transcriptomic signatures. These reference signatures are then used to decompose the observed mRNA counts in each Visium spot and infer the abundance of each cell type⁴⁵. The model takes the following as input: (1) the raw [$\text{spots} \times \text{genes}$] ST gene count matrix, (2) the annotated [$\text{cells} \times \text{genes}$] scRNA-seq gene count matrix that was labeled from “Bioinformatics preprocessing,” and (3) a prior distribution of the number of cells per spot (either uniform or spot-specific). The output is a posterior distribution for the number of cells of each type in each Visium spot in the domain.

In our analysis, all parameters were set to default except for the following: (1) the expected cell abundance per location, which was set to 20, in line with the approximate average number of nuclei per spot detected in the H&E images in the original experimental study¹⁸, and (2) the per-location normalization regularization parameter ($\alpha_{\text{detection}} = 20$), which was also adjusted in the original study to account for large variations in mRNA detection sensitivity across different Visium spots¹⁸. Furthermore, in this study, we will use the 5% quantile cell number of the posterior distribution to summarize the spatial distribution of cell types across the Visium spots. This choice reflects a conservative estimate that emphasizes high-confidence cell presence¹⁸ while mitigating the influence of outliers or uncertain assignments. Nonetheless, our computational pipeline can readily accommodate the use of the full posterior distribution of the cell numbers with some corresponding adjustments if desired; see Supplementary Note 7.

To infer the TGF β -mediated cell-cell interaction strengths from the transcriptomics data, we use CellPhoneDB³¹, a pioneering and widely used computational toolkit for reconstructing cell-cell communication networks across different cell types from scRNA-seq data. The toolkit leverages a comprehensive, curated, and annotated database of ligand-receptor interactions, including multi-subunit heteromeric complexes, that is compiled from public literature resources, to carry out a systemic analysis of cell-cell communication as motivated by the secreted diffusing ligands³². The ligand-receptor interaction strengths between one secreting cell type and another receiving cell type are derived on the basis of the mean mRNA co-expression of the receptor (or receptor subunits) in one cell type and the ligand in another; see ref. 46 for details on the specific calculation of interaction strengths. Note that these interactions are directional and not necessarily symmetric (i.e., interaction strength from cell type A to B is not the same as from B to A).

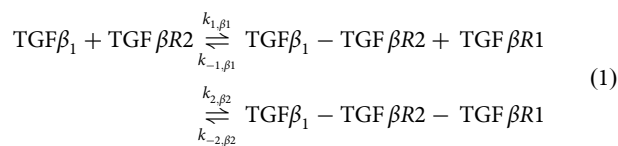
In our case, the CellPhoneDB algorithm would take the following as input: (1) the annotated (by cell type) cell-gene expression matrix obtained from processing the scRNA-seq data, (2) the DGE file that features the genetic signature for each cell type, which was used for the scRNA-seq clustering annotation, and (3) an annotated database of the relevant ligand-receptor interactions, for which we use the default CellPhoneDB database. To ensure biological relevance, only ligand-receptor pairs in which both components are expressed in more than 10% of cells within their respective clusters are considered. Furthermore, at least one component of the ligand-receptor pair must be listed in the DGE file for the interaction to be considered.

CellPhoneDB uses the co-expression of ligand-receptor pairs between cell types as a proxy for estimating interaction strength. While this approach provides a practical and interpretable measure and has its advantages, it also comes with limitations, which are examined in the “Discussion,” along with discussing other potential toolkits that can be used. Nonetheless, at this stage, we chose CellPhoneDB due to its widespread adoption, ease of integration, and suitability for proof-of-concept demonstrations of our computational pipeline^{6,21,31,32}. In Supplementary Note 1, we further illustrate the utility of our computational pipeline using another open-source high-

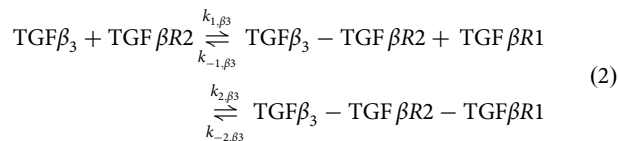
throughput sequencing dataset, derived from a skin tissue sample of Acne keloidalis, a chronic inflammatory condition primarily affecting the scalp and neck, characterized by papules, pustules, hair loss, and keloidal scarring²¹. In this example, interaction strengths are spatially inferred using a recently developed bioinformatics tool²⁷ that accounts for the spatial distances between cells. This demonstrates how our computational pipeline can readily accommodate alternative tools for deriving interaction strengths, depending on the specific needs and goals of the user. Importantly, our pipeline is modular and can be adapted to incorporate alternative toolkits, methods, or metrics for quantifying interaction strengths, depending on the context and specific needs (see “Discussion”).

TGFβ binding kinetics

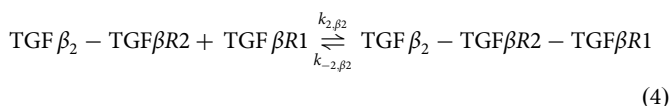
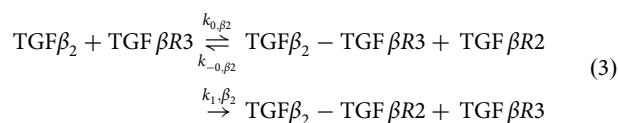
In humans, TGFβ₁, TGFβ₂, and TGFβ₃ are the three main isoforms of the transforming growth factor-beta (TGFβ) family. Each isoform initiates signaling by binding to specific surface receptors, leading to the formation of receptor-ligand complexes that in turn trigger intracellular cascades. TGFβ₁, the most prominent and well-studied isoform in wound healing, first binds to the type II receptor (TGFβR2). This complex then recruits the type I receptor (TGFβR1) to form a heterotrimeric complex capable of initiating intracellular SMAD signaling:



Likewise, TGFβ₃ follows a similar receptor-binding mechanism:



In contrast, TGFβ₂ has low affinity for TGFβR2 and typically requires the presence of the co-receptor TGFβR3 (also known as betaglycan), which facilitates ligand presentation to TGFβR2 before the aforementioned binding cascade proceeds:



At steady state, such as during the wound remodeling phase, the net flux through each step of these binding cascades can be assumed to be equal (i.e., the formation and consumption rates of each intermediate complex are balanced)⁴⁷. Therefore, to simplify the modeling of ligand-receptor dynamics within our reaction-diffusion framework, we focus only on the first binding event for each isoform, with corresponding binding rates ([Ms]⁻¹) for isoforms 1, 2, and 3 denoted as k_1 , k_2 , and k_3 , respectively. This allows us to represent ligand uptake via a single effective reaction step, while retaining the flexibility to extend to multistep kinetics in future refinements (see *Discussion*):



Considering the ligand-receptor interactions as defined in Eqs. (5)–(7), the resulting interaction strengths across the cell-type pairs and TGFβ isoforms as inferred from the transcriptomics data are shown in Fig. 3.

Mathematical model

The spatial distribution of a ligand (for example, a growth factor g) can be modeled via the following reaction-diffusion equation:

$$\nabla \cdot (\mathbf{D}_g \nabla g) - \lambda_g g + S_g(x) = \frac{\partial g}{\partial t} \quad (8)$$

where, in the steady-state case, $\frac{\partial g}{\partial t} = 0$. Here, $\mathbf{D}_g$ (in m²s⁻¹) denotes the diffusivity tensor of the growth factor g ; when assumed isotropic and spatially uniform, this simplifies to a scalar diffusion coefficient. The term λ_g (in s⁻¹) represents the background decay rate of g in the domain, and $S_g(x)$ accounts for spatially varying sources or sinks corresponding to production or absorption of the ligand. The source term $S_g(x)$ depends on the production by the different cell types $i \in \mathcal{C}$ where $\mathcal{C}$ denotes the set of all cell types, and on the local uptake of g due to binding to cell-surface receptors.

The choice between a continuum or particle-based formulation of Eq. (8) in the computational pipeline depends partly on the nature and resolution of the available data. Methods like standard Visium ST, which map gene expression to spatially localized spots encompassing multiple cells, and which can be combined with scRNA-seq to infer cell densities, as carried out in “Bioinformatics preprocessing,” would be better suited for a continuous model. In this section, we adopt a continuum-based formulation, in line with the nature of the data (i.e., the cell density distributions) that could be extracted in “Bioinformatics preprocessing.” In contrast, when spatially resolved single-cell data is available, for instance, from technologies such as Visium HD, Multiplex Immuno-fluorescence, or Fluorescence In Situ Hybridization, particle-based formulations would be more suitable. In Supplementary Note 2, we derive the particle-based formulation and outline how it can be incorporated into the pipeline when the precise cellular positions are available. Ultimately, the goal is to accommodate various spatial resolution data types while maintaining compatibility between the particle-based and continuum approaches.

In this study, we employ the FV method, which is well-suited for solving the type of conservation laws in Eq. (8). Of course, other numerical methods may also be utilized, provided they follow the same core steps; that is, they proceed from discretizing both the equation and spatial domain, leading to a set of algebraic equations defined over CVs or finite elements, which can then be assembled into a system to solve for the spatial distribution of g . For consistency in units and to accurately reflect the geometry of the experimental data, we formulate the FV equations in three dimensions. This is necessary because certain parameters—such as molecular binding rates—are defined in volumetric units (e.g., [Ms]⁻¹ = 10⁻³m³/(mol s)), requiring a 3D spatial context for proper dimensional alignment. Additionally, the tissue section under analysis, though visualized and discretized in 2D using a hexagonal grid, possesses a finite thickness in the third dimension. To account for this, we model each CV as a hexagonal prism, enabling us to incorporate the sample’s depth while preserving the natural hexagonal tiling used in ST platforms. Once the 3D formulation is established, we reduce the system to an effective 2D representation by collapsing the thickness dimension.

Using the divergence theorem on Eq. (8), and assuming that the diffusivity tensor is isotropic and uniform, the divergence of the field ($\mathbf{D}_g \nabla g$) inside the CV can be related to the flux of the vector field through the CV

closed surface boundary:

$$\int_{CV} \nabla \cdot (D_g \nabla g) \partial \mathbf{V} = \oint_B D_g \nabla g \cdot d\mathbf{S} = D_g \sum_{f \in \text{surfaces}} \left(\int \nabla g \cdot \hat{\mathbf{n}}_f ds_f \right) \quad (9)$$

where $\hat{\mathbf{n}}_f$ is the outward-directed unit normal vector to surface f . Given that Eq. (8) is purely diffusion-based, we utilize a collocated arrangement for the FV method in which the values for g are stored in the FV centroids. Furthermore, since the PDE in Eq. (8) is second-order in space, (at least) a second-order spatial discretization is warranted. Hence, the flux through the CVs can be formulated in terms of the value of ∇g at the edge centroids and can be written as:

$$D_g \sum_{f \in \text{surfaces}} \left(\int \nabla g \cdot \hat{\mathbf{n}}_f ds_f \right) \approx h D_g \sum_{e \in \text{edges}} ((\nabla g)^{f_e} \cdot \hat{\mathbf{n}}_e L_e) \quad (10)$$

where h is the thickness of the CV prism, L_e is the length of the CV edge e in 2D, and $(\nabla g)^{f_e}$ is the gradient of g evaluated at the centroid (f) of edge e . In a uniform orthogonal grid with no mesh non-orthogonality or skewness, the $(\nabla g)^{f_e}$ at the edge can be evaluated in terms of the value of g at the two centroids on the opposite sides of the edge, using the simple central difference scheme. This gives rise to $\phi_g^D(C)$ in Eq. (11), which denotes the total diffusive flux of the growth factor out of CV C , in terms of the values of g at the centroid of the C centroid (g_c) and those at the centroids of the neighboring CVs ($g_i, i \in N(C)$):

$$\phi_g^D(C) = h D_g L_r \left(\left(\sum_{i \in N(C)} g_i \right) - |N(C)| g_c \right) \quad (11)$$

where $L_r = L_e/L_C$ and L_C is the distance between the centroids in the uniform mesh, and $|N(C)|$ is the cardinal of $N(C)$ (i.e., number of CVs adjacent to C). $S_g(x)$ in Eq. (8) would represent any source or sink terms inside each CV due to ligand production or absorption by the different cell types, respectively. This can be written as:

$$\int_C S(x) d\mathbf{V} = hA \left(\sum_{i=1}^{|C|} (\rho_g^i d_c^i - k_g^i R_g^i d_c^i) \right) \quad (12)$$

where A is the area of the CV, ρ_g^i is the production rate (mol/(s.cell)) of ligand g by cells of type i , d_c^i is the density of the cells of type i in the CV C , R_g^i is the number of ligand g -binding receptors (mol/cell) on cell type i and k_g^i is the binding rate (Ms)⁻¹ per cell between growth factor g and its corresponding receptor on cell type i . The definition is set up such that in the $\lim_{A \rightarrow 0}$, we recover the particle-based point source formulation in Supplementary Note 2. Note that in case there are multiple receptor types j binding ligand g on cell type i , the corresponding absorption rate inside the CV simply becomes $d_c^i \sum_j k_g^{i,j} R_g^{i,j}$.

Similarly, the decay rate within each CV can be formulated as:

$$\int_C \lambda_g g d\mathbf{V} = hA \lambda_g g_c \quad (13)$$

Equations (10), (12), and (13) can be divided by the thickness h and combined to form of a linear balance equation that accounts for the conservation of g within each CV C surrounded by neighboring CVs $j \in N(C)$ in the uniform orthogonal FV mesh:

$$\left(D_g L_r |N_C| + A \left(\lambda_g + \sum_{i=1}^{|C|} R_g^i k_g^i d_c^i \right) \right) g_c - D_g L_r \sum_{j \in N(C)} g_j = A \left(\sum_{i=1}^{|C|} \rho_g^i d_c^i \right) \quad (14)$$

where the variables A , L_r and $|N_C|$ depend on the choice of the mesh and CV shape. Assuming no-flux Neumann boundary conditions, Eq. (14) can be formulated for each CV and assembled to form a linear system of equations that can be solved for the growth factor concentrations across all CVs $\mathbf{g}$:

$$\mathbf{A}(\Theta_g) \mathbf{g} = \mathbf{b}(\Theta_g) \quad (15)$$

The entry values for matrix $\mathbf{A}$ and the column vector $\mathbf{b}$ in Eq. (15) depend on the model parameter values, which are discussed in “Parameter initialization and assembly.” The concentration of the growth factors at the receptor sites can then be used to calculate the ligand-receptor interaction strengths for the different ligand-receptor pairs on the various cell types across the tissue domain of interest. In the case of the continuous model, the interaction strength between cell type l secreting ligand A and cell type r expressing receptor B and can be approximated as:

$$\mathcal{I}_{A-B}^{l-r} = \frac{1}{|CV|} \sum_{i=1}^{|CV|} \left(R_B^l k_B^r g_A^i \frac{\rho_A^l d_l^i}{\sum_{j \in C} \rho_A^j d_j^i} \right) \quad (16)$$

where $|CV|$ is the number of CVs in the domain, g_A^i is the concentration of growth factor A in the centroid of CV i , d_l^i is the cell density of cell type l in CV i , and C is the set of all cell types of interest. Equation (16) infers the contributions of a particular cell type in a particular spot to the total ligand binding in that spot.

Since the relative scaling and offset between the model-derived and experimental interaction strengths have yet to be established and likely depend on the specific bioinformatics tool used for inference, a correlation-based comparison, such as Pearson’s correlation coefficient r , can be employed to compare the two and inform parameter calibration. Such a measure is independent of absolute magnitudes and can identify the optimal model parameters that produce interaction strengths which, when compared pairwise, preserve the same relative ratios as those observed in the experimental data. While Pearson’s correlation is used here, alternative monotonic (but not necessarily linear) relationships may also be considered. Future work can explore the precise relationship between model-derived interaction strengths and those inferred by specific bioinformatics pipelines (see “Discussion”). Furthermore, in this section, we demonstrate how, under certain assumptions, namely that of additive, independent, and homoscedastic Gaussian noise in the interaction strengths, which are revisited in “Discussion” and Supplementary Note 7, minimizing the likelihood function to infer the optimal parameter combination becomes equivalent to maximizing the correlation between the data and model-derived interaction strengths.

Likelihood-based analysis

We assume a linear relationship between the transcriptomics-derived and model-derived interaction strengths, expressed as:

$$\mathcal{I}_i^{\text{data}} = s \mathcal{I}_i^{\text{model}} + b + \epsilon_i \quad (17)$$

where s and b are the universal (i.e., shared across the data points) scaling factor and offset parameters, respectively, and ϵ_i is the uncertainty, or noise term associated with interaction strength $\mathcal{I}_i$ for data point $i, i = 1 \dots N$.

We shall refer to the parameters s , b , and the noise covariance matrix $\Sigma = f(\sigma_i), i = 1 \dots N$ in Eq. (17) as the static parameters, to differentiate them from the original parameters θ to be calibrated, henceforth called dynamic parameters (since they affect the results of the simulated states). The static parameters are often unknown and thus have to be estimated along with the dynamic parameters. To infer the static and dynamic unknown parameters, we maximize the likelihood of observing the transcriptomics-based interaction strengths given the parameters. If we assume a multivariate normal

distribution of the data residuals (noise terms):

$$\mathcal{I}^{\text{data}} \sim \mathcal{N}(s\mathcal{I}^{\text{model}} + b, \Sigma) \quad (18)$$

Then the likelihood function can be formulated as

$$L(\theta, s, b, \Sigma) = \frac{1}{(2\pi)^{N/2} |\Sigma|^{1/2}} \exp\left(-\frac{1}{2} (\mathcal{I}^{\text{data}} - (s\mathcal{I}^{\text{model}} + b))^T \Sigma^{-1} (\mathcal{I}^{\text{data}} - (s\mathcal{I}^{\text{model}} + b))\right) \quad (19)$$

Instead of maximizing L directly, it is customary and numerically often preferable to minimize the NLL, denoted here as $J(\theta, s, b, \Sigma)$:

$$J(\theta, s, b, \Sigma) = \frac{1}{2} (\mathcal{I}^{\text{data}} - (s\mathcal{I}^{\text{model}} + b))^T \Sigma^{-1} (\mathcal{I}^{\text{data}} - (s\mathcal{I}^{\text{model}} + b)) + f(\Sigma) \quad (20)$$

If we assume that the noise terms are independent, then the likelihood function gives rise to the conventional product-based expression:

$$L(\theta, s, b, \sigma) = \prod_i (\mathcal{I}_i^{\text{data}} | s\mathcal{I}_i^{\text{model}} + b, \sigma_i) \quad (21)$$

where σ is the vector of the data point variances $\sigma_i, i = 1 \dots N$. Subsequently, the corresponding NLL simplifies to the following expression for Gaussian independent noise:

$$J(\theta, s, b, \sigma) = \frac{1}{2} \sum_i \left(\log(2\pi\sigma_i^2) + \frac{(\mathcal{I}_i^{\text{data}} - (s\mathcal{I}_i^{\text{model}} + b))^2}{\sigma_i^2} \right) \quad (22)$$

The standard approach to minimizing Eq. (22) is to consider the extended parameter vector $[\theta, s, b, \sigma]$ and optimize all its elements simultaneously, which can make the optimization problem significantly harder to solve⁴⁸. Alternatively, in the hierarchical scheme introduced by Schmiester et al.⁴⁸, the optimization problem is split into an outer problem where the dynamic parameters θ are optimized, and an inner problem where the static parameters are optimized conditioned on θ . In other words, we compute:

$$\arg \min_{\theta} J(\theta) \text{ with } J(\hat{\theta}) = J(\theta, s(\theta), b(\theta), \sigma) \quad (23)$$

where the optimization of the inner problem is given by:

$$[s(\theta), b(\theta)] = \arg \min_{s,b} J(s, b | \theta, \sigma) \quad (24)$$

It has been shown that the hierarchical splitting exhibited in Eqs. (23) and (24) preserves the global optima of the standard optimization problem⁴⁸. Normally, the inner optimization problem in Eq. (24) needs to be solved numerically. However, given that the static parameters s and b are shared among the data points, if the noise variance parameters are also shared (i.e., the homoscedastic assumption $\sigma_i = \sigma$), then an analytic condition for the optimal static parameters can be derived by solving for:

$$\nabla_{s,b} J(\theta, s, b, \sigma) = 0 \quad (25)$$

where the optimal solutions are given by⁴⁸:

$$\hat{s}(\theta) = \left(\sum_{i=1}^N (\mathcal{I}_i^{\text{model}})^2 \right)^{-1} \sum_{i=1}^N ((\mathcal{I}_i^{\text{data}} - b) \mathcal{I}_i^{\text{model}}) \quad (26)$$

$$\hat{b}(\theta) = \frac{1}{N} \sum_{i=1}^N (\mathcal{I}_i^{\text{data}} - s\mathcal{I}_i^{\text{model}}) \quad (27)$$

Inserting Eq. (27) into Eq. (26) and solving the system of equations, it can be shown that the optimal solution for the variables can be given by:

$$\hat{s}(\theta) = \frac{\text{cov}(\mathcal{I}^{\text{data}}, \mathcal{I}^{\text{model}})}{\text{Var}(\mathcal{I}^{\text{model}})} \quad (28)$$

$$\hat{b}(\theta) = \mathbb{E}[\mathcal{I}^{\text{data}}] - s\mathbb{E}[\mathcal{I}^{\text{model}}] \quad (29)$$

Note that the optimal parameters in Eqs. (28) and (29) correspond to the regression solution of minimizing the sum of the squared errors between the model and data interaction strengths, as expected. Furthermore, since it can be shown that the second derivatives are positive⁴⁸, then the solution is a global minima. Note that the solution in Eqs. (28) and (29) assumes that all data points share the same value for the noise variance σ , which may be specified or optimized as a static parameter, depending on the problem setup⁴⁸; in our model, they are treated as fixed (for instance, inferred from the experimental noise and computational propagated uncertainty; see Supplementary Note 7). Next, we insert the analytical solution to our inner optimization problem in Eqs. (28) and (29) into the outer optimization problem in Eq. (23). Inserting Eq. (29) into the NLL expression of Eq. (22) we obtain:

$$J(\theta) = N \log(\sqrt{2\pi}\sigma) + \frac{1}{2\sigma^2} \sum_{i=1}^N ((\mathcal{I}_i^{\text{data}} - (\hat{s}(\theta)\mathcal{I}_i^{\text{model}}(\theta) + \mathbb{E}[\mathcal{I}^{\text{data}}] - \hat{s}(\theta)\mathbb{E}[\mathcal{I}^{\text{model}}(\theta)]))^2 \quad (30)$$

$$= N \log(\sqrt{2\pi}\sigma) + \frac{1}{2\sigma^2} \sum_{i=1}^N ((\mathcal{I}_i^{\text{data}} - \mathbb{E}[\mathcal{I}^{\text{data}}]) + \hat{s}(\theta)(\mathcal{I}_i^{\text{model}}(\theta) - \mathbb{E}[\mathcal{I}^{\text{model}}(\theta)]))^2 \quad (31)$$

Expanding Eq. (31) and inserting the analytical solution for $\hat{s}(\theta)$ from Eq. (28) into the expression, as well as the definitions of variance and covariance between random variable(s):

$$\text{var}(X) = \frac{1}{N} \sum_{i=1}^N (x_i - \mathbb{E}[X])^2 \quad (32)$$

$$\text{cov}(X, Y) = \frac{1}{N} \sum_{i=1}^N (x_i - \mathbb{E}[X])(y_i - \mathbb{E}[Y]) \quad (33)$$

We obtain an expression for the likelihood function in terms of the data and model summary statistics:

$$J(\theta) = N \log(2\pi\sigma) + \frac{N}{2\sigma^2} \left(\text{var}(\mathcal{I}^{\text{data}}) - \frac{\text{cov}^2(\mathcal{I}^{\text{data}}, \mathcal{I}_i^{\text{model}}(\theta))}{\text{var}(\mathcal{I}_i^{\text{model}}(\theta))} \right) \quad (34)$$

Finally, rearranging Eq. (34) and noting the definition of Pearson's correlation coefficient (r):

$$r(\mathcal{I}^{\text{data}}, \mathcal{I}_i^{\text{model}}(\theta)) = \frac{\text{cov}(\mathcal{I}^{\text{data}}, \mathcal{I}_i^{\text{model}}(\theta))}{\sqrt{\text{var}(\mathcal{I}^{\text{data}})\text{var}(\mathcal{I}_i^{\text{model}}(\theta))}} \quad (35)$$

we obtain an expression for the NLL in terms of r :

$$J(\theta) = N \log(2\pi\sigma) + \frac{N\text{var}(\mathcal{I}^{\text{data}})}{2\sigma^2} (1 - r^2(\theta)) \quad (36)$$

where $r(\theta)$ is defined in Eq. (35). In summary, we have shown that under the aforementioned assumptions, minimizing the NLL is equivalent to

maximizing $r(\theta)$:

$$\hat{\theta} = \arg \min_{\theta} (1 - r^2(\theta)) = \arg \max_{\theta} r(\theta), r \in [0, 1] \quad (37)$$

One advantage of adopting a likelihood-based formulation for parameter inference is that, when feasible, it enables us to construct confidence intervals around our MLEs; see Supplementary Note 3. These intervals can provide valuable insights into the reliability of our estimators by quantifying uncertainty. This is achieved through Fisher information theory, which characterizes how sensitive the data distribution is to changes in the parameters. Intuitively, higher Fisher information means the data change more noticeably with parameter variations, allowing for more precise parameter estimation and narrower confidence intervals. Using the Fisher information, the confidence intervals can be approximated as (see Supplementary Note 3):

$$\text{CI}_{100(1-\alpha)\%}(\hat{\theta}_i) \approx 2\sigma_{\alpha/2} \sqrt{[G^T(\hat{\theta})G(\hat{\theta})]^{-1}_{ii}} \quad (38)$$

where σ^2 is the homoscedastic noise variance, $z_{\alpha/2}$ is the quantile of the standard normal (e.g., 1.96 for 95% CI), and $G(\hat{\theta})$ is the Hessian of the NLL function with respect to the parameters. These relative confidence intervals are computed and used to provide important insights into both parameter identifiability and estimator reliability; see Supplementary Notes 4 and 5.

Parameter initialization and assembly

The first step in the parameter calibration framework involves constructing a prior distribution from which an ensemble of realizations can be initialized. To achieve this, we survey the current literature to either directly extract relevant parameter values or extrapolate them from available experimental and computational studies. Table 1 summarizes the key extracted parameters and outlines the rationale behind the assigned values and distribution.

The LogUniform distribution used for some parameters in Table 1 yields a right-skewed profile that inherently favors smaller values. This choice is appropriate when literature-reported values tend to cluster at the lower end of the range, while still accommodating the full span of observed magnitudes with greater weighting on the more frequently reported smaller values.

Quantitative secretion rates of TGF β per cell are not directly reported for macrophages or endothelial cells. However, M2 macrophages are well recognized as significant TGF β producers during wound healing, whereas endothelial cells (ECs) contribute to a lesser extent, particularly in the remodeling phase. Accordingly, as an initial estimate, the production rate distributions for macrophages and ECs were scaled to 50% and 20% of the fibroblast baseline distribution in Table 1, consistent with the functional hierarchy $\rho_{\text{fib}} \geq \rho_{\text{mac}} \geq \rho_{\text{EC}}$ described by Deng et al.¹⁹

In our analysis, each TGF β isoform is treated as an independent ligand whose spatial distribution is governed by the reaction-diffusion equation described in “Mathematical model,” giving rise to a concentration field for each isoform. The concentration fields are combined with selected model parameters to compute the interaction strengths across the cell types for each isoform, according to Eq. (16). The interaction strengths for each isoform can then be expressed in the form of an $n \times n$ matrix (where n is the number of the cell types of interest); this format mirrors the structure of experimentally inferred interaction strengths derived from high-throughput genetic sequencing data, as illustrated in Fig. 3.

There are currently no pharmacological tools available to easily distinguish between TGF β receptor subtypes, either via ligand-binding assays or through receptor-specific agonists or antagonists¹⁹. Consequently, the distribution of receptor subtypes cannot be readily determined. Likewise, the reported TGF β production rates do not differentiate among the three isoforms of the ligand. To account for this, our model randomly partitions the total reported TGF β production rates and receptor abundances for each

cell type among the three isoforms, introducing additional degrees of freedom to the parameter space. The parameters are then grouped into three distinct sets, one per isoform. For each realization, the diffusion coefficient (D) and degradation rate (λ_g) are assumed to be shared across all isoforms, reflecting the high similarity in molecular profiles among the TGF β isoforms. This assumption helps avoid unnecessary expansion of the parameter space, thereby reducing the risk of overfitting. Note that, in the specific case of analyzing data from normal skin wound healing at day 30 post-incision¹⁸, three distinct fibroblast clusters were identified. Each cluster was treated as a separate cell type, with the corresponding parameter values sampled from the fibroblast-specific distributions listed in Table 1.

Using each isoform-specific parameter set, the growth factor concentration field is computed by solving the FV linear system described in Eq. (15). The resulting spatial distributions are then used to calculate the corresponding interaction strength matrices between cell types according to Eq. (16). Finally, the interaction strengths across all isoforms and cell types are concatenated, flattened, and compared against bioinformatics-derived interaction strength data. The Pearson correlation coefficient is computed to quantify the agreement between model-predicted and data-inferred interaction patterns. The initialization, partitioning, and reassembly steps used in this parameter assembly framework are illustrated in Fig. 14.

Parameter calibration schemes

In this section, we outline the development of a parameter calibration framework that integrates the bioinformatics-derived data on the spatial distribution of cell densities and interaction strengths (see “Bioinformatics preprocessing”) with the reaction-diffusion numerical modeling approach described in “Mathematical model.” The first step involves reconstructing the spatial in silico discretization mesh used in the FV numerical scheme, based on the spatial positions of the Visium spots arranged on a hexagonal grid. Given that the resulting computational mesh is uniform, orthogonal, and non-skewed, it can be fully characterized using two key pieces of information extracted from the Visium ST layout.

First, the geometric properties of the mesh must be extracted; in particular, the CV area A and the ratio of the edge length to the center-to-center distance L_r used in Eq. (14) can be obtained from the knowledge of the hexagon CV diameter (84.5 μm). Second, the connectivity between adjacent CVs must be determined. This can be captured by a square Boolean adjacency matrix where each i, j entry indicates whether CVs C_i and C_j are directly adjacent, and where adjacency is defined via the indicator function $\mathcal{I}(|C_i - C_j| \leq L_c + \epsilon)$, with L_c representing the center-to-center distance between neighboring hexagonal CVs and ϵ a small numerical tolerance.

These mesh-related components, along with the cell density maps for each cell type and the experimentally inferred interaction strength matrices, are treated as fixed inputs throughout the parameter calibration process. In contrast, the model parameters listed in Table 1 and shown in Fig. 14 are treated as variable parameters subject to calibration. Below, we discuss several relevant parameter calibration schemes, including gradient-based optimization schemes and ABC schemes, highlighting the advantages and limitations of each, and describing how they can be sequentially combined within a hierarchical calibration framework that leverages the strengths of each approach while mitigating their respective weaknesses.

The core principle behind using gradient-based optimization schemes is that the inputs (i.e., parameter realizations) are passed through a sequence of well-defined, differentiable mathematical operations to compute the final metric; in this case, it is the correlation-based objective function in Eq. (37). This structure enables the use of back propagation to trace how changes in the inputs affect the output, allowing us to iteratively update the parameters and minimize the discrepancy between the model-predicted and data-inferred interaction strengths.

To prevent overfitting and to constrain the parameters within biologically plausible ranges, we incorporate a regularization term in the loss function. For each parameter $i \in [1, d]$ in the standardized parameter space, where d is the total number of tunable parameters, we define an acceptable range $R_i = [a_i, b_i]$. These bounds are informed by, or slightly broader than,

Table 1 | Initial parameter values for the TGFβ reaction-diffusion model

Parameter [units]	Values [source(s)]	Initialization	Notes
D [$\times 10^{-11} \frac{\text{m}^2}{\text{s}}$]	1. 2.6 [58] 2. 2.13 [59] 3. 13 [60] 4. 2.9 [10]	LogUniform ([1, 20])	- The same D is assumed across all isoforms since they share similar molecular weights (25 kDa) and (80%) sequence identities. - The second value reflects the ligand's interaction with the cartilage extracellular matrix (ECM), similar to what is observed in dermal tissues, leading to slower diffusion compared to free diffusion⁶⁰.
λ_g [$\times 10^{-6} \text{s}^{-1}$]	1. 13.32 [61←62] 2. 4.10 [10←39]	Uniform ([2, 20])	- The same λ_g is assumed across all isoforms since they share similar molecular weights (25 kDa) and (80%) sequence identities. - The first value is based on a TGFβ life cycle of 48 h used in the computational model⁶¹ based on the experimental results⁶²; we assume that this corresponds to a 90% decay of initial concentration in order to calculate the decay constant value⁶⁰.
k_1 [$\times 10^2 \frac{\text{m}^3}{\text{s} \cdot \text{mol}}$]	1. 5.1–5.7 [14,15] 2. 7.2 [13,17] 3. 11.6 [13←16] 4. 210–310 [15]	LogUniform ([2, 500])	The larger range of values from the last source corresponds to the experimental conditions in which the receptors were immobilized ¹⁵ .
k_2 [$\times 10^2 \frac{\text{m}^3}{\text{s} \cdot \text{mol}}$]	15 [13←63]	LogUniform ([5, 200])	
k_3 [$\times 10^2 \frac{\text{m}^3}{\text{s} \cdot \text{mol}}$]	1. 7–7.8 [17] 2. 6.9–7.9 [64] 3. 7.4 [13←17] 4. 18 [13←16] 5. 730–830 [15]	LogUniform ([5, 1000])	The larger range of values from the last source corresponds to the experimental conditions in which the receptors were immobilized ¹⁵ .
ρ_{fib} [$\times 10^{-22} \frac{\text{mol}}{\text{s} \cdot \text{cell}}$]	1. 1.556 [65] 2. 3.472 [65]	Uniform ([1.0, 4.0])	- The first value is from liver cells in a microfluidic chip, and converted to moles assuming a TGFβ molecular weight of 25 kDa. - The second value is based on the reported production rates in scleroderma and normal dermal fibroblast cells.
ρ_{mac} [$\times 10^{-22} \frac{\text{mol}}{\text{s} \cdot \text{cell}}$]	N/A	$0.5 \times \text{Uniform}([1.0, 4.0])$	Assumed to have a uniform distribution over half the range of ρ_{fib} , scaled by 0.5, since TGFβ secretion by fibroblasts is generally greater than that by macrophages during the remodeling phase of wound healing ⁴⁹ .
ρ_{EC} [$\times 10^{-22} \frac{\text{mol}}{\text{s} \cdot \text{cell}}$]	N/A	$0.2 \times \text{Uniform}([1.0, 4.0])$	Assumed to have uniform distribution over one fifth the range of ρ_{fib} , scaled by 0.2, since TGFβ secretion by ECs is generally the least compared to that by fibroblasts and macrophages during the remodeling phase of skin wound healing ⁴⁹ .
R_{fib} [$\times 10^{-20} \frac{\text{mol}}{\text{cell}}$]	1.387–2.467 [65]	Uniform ([1.32, 2.50])	Based on the reported number of receptors in scleroderma and normal dermal fibroblasts.
R_{mac} [$\times 10^{-20} \frac{\text{mol}}{\text{cell}}$]	0.0631–0.0640 [66,67] (for monocytes)	LogUniform ([0.049, 2.33])	The reported values correspond to monocytes; upon activation to macrophages, receptor numbers can rise significantly. This justifies the chosen starting distribution, which spans the lower monocyte values up to levels typical of fibroblasts and endothelial cells.
R_{EC} [$\times 10^{-20} \frac{\text{mol}}{\text{cell}}$]	0.997–2.193 [68]	Uniform ([0.990, 2.32])	The range of values covers the number of TGFβ receptors under physiological and receptor expression-inhibited conditions.

The values were collected from the literature and converted to SI units for dimensional consistency. This includes the cellular density of TGFβ receptors, which was converted from receptor count to number of moles by dividing by Avogadro's constant. The notation $\leftarrow$ indicates that the value(s) was extracted from the computational study in A based on the experimental results in B. k_1 , k_2 , and k_3 correspond to the binding rates of TGFβ isoforms 1, 2, and 3 to their corresponding TGFβ type 2, 3, and 2 receptors, respectively. The LogUniform distribution of a variable x in the range $[a, b]$ corresponds to uniform sampling in $\log_{10}$ space such that $\log_{10}(x) \sim U[\log(a), \log(b)]$.

the prior parameter ranges reported in Table 1. We then define a regularization function L_R that penalizes parameters falling outside their respective allowed intervals:

$$L_R(\theta) = \sum_{i=1}^d \text{RELU}(\max(\theta_i - b_i, a_i - \theta_i)) \quad (39)$$

The total loss function is given by a weighted sum of the fitting loss and the regularization loss, with a hyperparameter λ controlling the trade-off between fitting accuracy and regularization:

$$L = (1 - r^2) + \lambda L_R \quad (40)$$

Note that this λ is distinct from the decay rate of the ligand λ_g . Figure 15 demonstrates how the input parameters (with a simplified representation using only two cell types) are propagated through a sequence of mathematical operators to obtain the final loss function.

Next, applying the chain rule, we can map the dependence of the loss function L on any parameter Θ_i through the intermediate interaction

strengths describing the crosstalk between the secreting cell type l and the receiving cell type r , across all TGFβ isoforms $j = 1, 2, 3$:

$$\frac{dL}{d\theta^i} = \sum_{j=1}^3 \sum_{(l,r) \in \mathcal{LR}} \left(\frac{\partial L}{\partial \mathcal{I}_j^{l-r}} \left(\frac{\partial \mathcal{I}_j^{l-r}}{\partial \theta^i} + \frac{\partial \mathcal{I}_j^{l-r}}{\partial \mathbf{g}_j} \cdot \frac{\partial \mathbf{g}_j}{\partial \theta^i} \right) \right) \quad (41)$$

Note that in the case of the TGFβ model, as only the parameters D and λ_g are shared across the isoforms, the partial derivatives of the loss with respect to the remaining parameters in Eq. (41) are relevant to one isoform and do not require summing over all isoforms. Also note that the growth factor concentration field for each j ligand (isoform) $\mathbf{g}_j$ is obtained from solving the linear system of equations in (15). In this system, both the LHS matrix $\mathbf{A}_j(\Theta)$ and the RHS vector $\mathbf{b}_j(\Theta)$ depend on the parameters Θ , implying that $\mathbf{g}_j$ carries an implicit dependence on the parameters as well. To compute the Jacobian of $\mathbf{g}_j$ with respect to the parameters Θ , we can implicitly differentiate Eq. (15):

$$\frac{\partial \mathbf{A}(\Theta)}{\partial \theta^i} \mathbf{g}_j + \mathbf{A}(\Theta) \frac{\partial \mathbf{g}_j}{\partial \theta^i} = \frac{\partial \mathbf{b}(\Theta)}{\partial \theta^i} \quad (42)$$

Rearranging the terms, we obtain a linear system to solve for the parameter derivative of the concentration field $\frac{\partial \mathbf{g}_i}{\partial \theta_j}$:

$$\mathbf{A}(\Theta) \frac{\partial \mathbf{g}_i}{\partial \theta_j} = \left(\frac{\partial \mathbf{b}(\Theta)}{\partial \theta_j} - \frac{\partial \mathbf{A}(\Theta)}{\partial \theta_j} \mathbf{g}_i \right) \quad (43)$$

Note that both the original system of equations in (15) and the Jacobian system in (43) involve the construction of the same matrix $\mathbf{A}(\Theta)$, which can be leveraged for computational efficiency to solve for both systems of equations at every iteration of the gradient-based algorithms by reusing the LU factorization of $\mathbf{A}(\Theta)$. Implementation-wise, the Jacobian matrix $\partial \mathbf{g}_i / \partial \Theta$ can be computed via a custom autograd function in Python, and then incorporated into back propagation chain of Eq. (41) which in turn can be implemented via a tensor and automatic-differentiation library like PyTorch⁵⁰.

Next, we outline the key advantages and limitations of applying gradient-based optimization schemes, particularly in the context of parameter calibration for reaction-diffusion models. As previously discussed, one major advantage of this approach lies in the fact that the entire computational pipeline, from input parameters to the final loss metric, is explicitly defined through a sequence of differentiable operations. This enables the precise back propagation of gradients, allowing the effects of changes in the input parameters to be directly mapped onto the loss function, and thereby facilitating targeted fine-tuning of the input parameters. In addition, the inclusion of a regularization term helps constrain the parameter values within biologically plausible ranges, reducing the risk of overfitting.

However, the optimization landscape in such models is likely to be non-convex and may contain multiple local minima. As a result, convergence to a global minimum is not guaranteed. This issue can be partially mitigated by initializing multiple optimization runs, or “seeds,” at different starting points that span the parameter space. However, each gradient calculation run can be computationally expensive, and increasing the number of seeds further compounds the computational cost. Furthermore, some gradient-based optimization algorithms, such as ADAM or SGD, require explicit learning rate designation (adapted or global). In a multi-dimensional parameter space, such as the TGF β model, different parameters exhibit varying degrees of sensitivity or constraint: some parameters must lie within narrow bounds to yield a low fitting loss, while others permit broader variation. Ideally, the relative learning rates across the parameters should be adapted to reflect their respective relative degrees of constraint. Other gradient-based algorithms automatically infer the learning rates: in a Newton method with a Hessian computation, parameters are updated using the inverse of the Hessian (second derivative matrix), which acts as a scaling matrix, inherently adapting step sizes per parameter based on the local curvature of the loss surface. However, if the Hessian is singular or near-singular, this could lead to numerical instability and divergence during the iteration.

Considering these factors, an effective strategy would be to carefully select a small number of well-informed initialization seeds that are likely to be near regions of low loss and that collectively provide good coverage of the parameter space. The optimization performance can also be improved by assigning parameter-specific relative learning rates (if necessary) that reflect the relative degrees of constraint or sensitivity among the parameters.

Next, we discuss the employment of ABC techniques to calibrate the ligand reaction-diffusion model parameters. ABC comprises a family of Bayesian inference methods that estimate posterior parameter distributions without requiring explicit evaluation of the likelihood function, a task that can be computationally intractable or prohibitively expensive for complex models⁵¹. This is particularly relevant in our setting, where evaluating and working with the full noise structure would make likelihood-based inference much more difficult. Let Θ be a parameter vector to be estimated. Given some prior distribution $\pi(\Theta)$, the goal is to approximate the posterior distribution that integrates knowledge from the available experimental data: $\pi(\Theta|x) \propto f(x|\Theta)\pi(\Theta)$, where $f(x|\Theta)$ is the likelihood of obtaining the data x

given the parameter space Θ ⁵¹. In ABC, the likelihood is replaced by a comparison between simulated and experimental data, using Monte Carlo simulations to explore parameter space. In our application, the comparison is based on model-predicted interaction strengths derived from the simulated ligand concentration fields. The final posterior distribution thus reflects the combined information from model dynamics, prior knowledge, and experimental measurements.

The simplest ABC algorithm would be the ABC rejection sampler⁵², the general steps for which, in the context of calibrating the reaction-diffusion model, are summarized below:

1. Sample candidate parameter vector Θ^* from some proposal (prior) distribution $\pi(\Theta)$. In our case, the prior distribution can be constructed by screening the experimental and computational data from the relevant literature as described in *Parameter Initialization and Assembly*.
2. Use the parameter vector Θ^* to simulate a dataset x^* from the dynamic model. In our case, this would correspond to the concentration fields of the ligand(s). The dataset x^* can be used to compute some secondary output y^* , such as the interaction strengths.
3. Compare the simulated output y^* with the corresponding experimental data (x_0 or y_0), using a scoring function or similarity measure and threshold ϵ . We use Pearson's correlation coefficient r as our similarity measure and accept Θ^* if $r(y_0, y^*) \geq \epsilon$, otherwise reject.
4. Return to step (1) and repeat for a predetermined number of trials or until the desired number of samples is achieved.

The model outputs are thus partitioned into two sets of accepted and discarded realizations, and the accepted set is mapped back into the input parameters⁵³ to provide an empirical approximation to the posterior distribution $f(\Theta|x)$.

The ABC rejection sampler offers several advantages. It is straightforward to implement computationally and, because steps 1–3 of the algorithm are independent, it is embarrassingly parallelizable⁵¹, for instance, across multiple cores within a high-performance computing (HPC) cluster. Moreover, by comparing the prior and posterior parameter distributions, one can gain valuable insight, at no additional computational cost, into parameter identifiability, the degree of constraint imposed by the data, and the sensitivity of model outputs to parameter variations; see ref. 51 for an example.

Unlike gradient-based optimizers, the ABC rejection sampler does not provide a mechanism to map discrepancies between model outputs and data-derived outputs back onto the inputs for precise parameter fine-tuning. Furthermore, the acceptance rate of the ABC rejection sampler can be very low when prior distributions are poorly informed or substantially different from the posterior distributions^{51,54,55}. In our model, while some priors can be constructed on physiological grounds from experimental data (see “Parameter initialization and assembly”), reliable prior approximations are not available for other parameters (see discussion in “Introduction”). This issue is exacerbated in multi-dimensional parameter spaces, where inefficiencies can arise due to mismatches between initial sampling and the target posterior distribution defined by the acceptance criterion $r(y_0, y^*) \geq \epsilon$. For example, a parameter may individually exhibit a high likelihood, but if combined with low-likelihood values of other parameters, the low joint posterior probability will result in a poor acceptance criterion of the former³⁰. Overall, when prior knowledge is limited for some parameters, the ABC rejection sampler can be viewed as a preliminary step that produces an improved starting realization of seeds from the prior distribution, which are better aligned with the ex vivo transcriptomics data, while simultaneously providing useful insight into parameter constraints and identifiability.

When the posterior distribution is far from the prior, sampling directly from the prior in the ABC rejection sampler is often inefficient due to low acceptance rates⁵⁴. To address this, Markov Chain Monte Carlo (MCMC) methods such as the Metropolis-Hastings algorithm can be employed. MCMC constructs a Markov chain whose stationary distribution corresponds to the target posterior $\pi(\Theta|x)$. For likelihood-free simulations, Marjoram et al.⁵⁴ introduced an MCMC scheme for generating observations

from a posterior distribution without the need for the likelihood calculation, which we refer to here as ABC MCMC. The ABC MCMC scheme proceeds as follows^{30,54}:

1. Initialize Θ_i , $i = 1$.
2. Generate a candidate parameter realization $\Theta^* \sim q(\Theta|\Theta_i)$ where q is the proposal density, or transition kernel. Common transition kernels include Gaussian distributions with a mean of Θ_i and some specified variance.
3. Generate the dataset x^* from the model and Θ^* and use it to compute the model output y^* .
4. Compare the simulated output, y^* , with the experimentally derived output y_0 using a similarity measure (r) and an acceptance threshold ϵ .
5. If $r(y_0, y^*) < \epsilon$, go back to step 2, otherwise set and accept $\Theta_{i+1} = \Theta^*$ with probability:

$$\alpha = \min\left(1, \frac{\pi(\Theta^*)q(\Theta_i|\Theta^*)}{\pi(\Theta_i)q(\Theta^*|\Theta_i)}\right) \quad (44)$$

otherwise set $\Theta_{i+1} = \Theta_i$.

6. if $i < N$ where N is the prescribed Markov Chain length, increment $i = i + 1$ and go to step 2.

Thus, ABC-MCMC algorithms generate a sequence of serially and highly correlated samples⁵⁴ which meet the criteria $r(y_0, y^*) \geq \epsilon$ after convergence is achieved. The major difference between the conventional and the ABC MCMC methods is that the former requires an explicit definition of the ratio of the likelihood functions to calculate the transition probability upon transition from Θ to Θ^* :

$$\alpha = \min\left(1, \frac{\pi(x|\Theta^*)}{\pi(x|\Theta_i)} \frac{\pi(\Theta^*)q(\Theta_i|\Theta^*)}{\pi(\Theta_i)q(\Theta^*|\Theta_i)}\right) \quad (45)$$

while in the latter case, if both Θ and Θ^* result in successful realizations ($r(y_0, y^*) \geq \epsilon$), and, assuming that the perturbation in the parameter values is small, the likelihood ratio is approximated as 1.

When the prior distribution is considerably different than the posterior distribution, the ABC-MCMC algorithm can deliver substantial increases in acceptance rate over ABC rejection sampler algorithm, as demonstrated in the case studies by Majoram et al.⁵⁴. One disadvantage of such MCMC algorithms, however, is that the resulting Markov chain samples are serially correlated, which, combined with potentially low acceptance probabilities or poorly chosen proposal distributions, can result in long chains that mix slowly or become trapped in low-probability regions^{30,51,55}. Furthermore, Metropolis-Hastings is known to scale poorly with the number of parameters to infer.

In theory, generating a large number of uncorrelated and independent realizations (or particles) can help prevent Markov chains from becoming trapped in low-probability regions. Moreover, by starting from a broad prior distribution with a relatively loose error tolerance and gradually tightening the tolerance, ABC algorithms can more effectively explore high-dimensional parameter spaces without getting stuck in realizations of local extrema probabilities⁵¹.

To this end, SMC ABC methods were introduced by Sisson et al.⁵⁵ as an improvement over the simple ABC rejection sampler. In SMC-ABC, a sequence of particle populations is generated, each evolving through intermediate distributions that transition smoothly from the prior toward the target posterior. Examples of this progressive narrowing of posterior distributions in practice can be found in refs. 30,56.

At each iteration t , the acceptance threshold ϵ_t is smaller than in the previous iteration ($\epsilon_t < \epsilon_{t-1}$), ensuring that the region of parameter space explored at iteration t is a subset of that explored at $t - 1$. Conceptually, this can be viewed as a sequence of rejection-sampling steps where the posterior from one iteration serves as the prior for the next. The gradual tightening of ϵ_t promotes efficient convergence to the posterior distribution^{30,51,55}. Starting

from a prior distribution $\pi(\Theta)$, the general ABC-SMC algorithm proceeds as follows⁵¹:

1. Initialize $\epsilon_1, \dots, \epsilon_T$. Set the population indicator $t = 0$.
2. Set the particle indicator $i = 1$.
3. Sample Θ^* from the previous population $\{\Theta_{t-1}^i\}$ with weights w^{t-1} and perturb the particle (according to some perturbation Kernel) to obtain $\Theta^{**} \sim K_t(\Theta|\Theta^*)$. If $\pi(\Theta^{**}) = 0$, repeat the sampling. If $t = 0$, then sample Θ^{**} directly from the prior $\pi(\Theta)$.
4. Simulate a candidate dataset (e.g., the ligand concentration field) from the obtained parameter realizations $x^* \sim f(x|\Theta^{**})$, calculate the output y^* (e.g., the interaction strengths), and calculate the corresponding similarity statistic ($r(y, y^*)$); if $r(y, y^*) < \epsilon_t$, return to step 3.
5. Set $\Theta_t^i = \Theta^{**}$ and calculate the weight for particle $\Theta_t^{(i)}$:

$$w_t^{(i)} = \begin{cases} 1 & \text{if } t = 0, \\ \frac{\pi(\Theta_t^{(i)})}{\sum_{j=1}^N K_t(\Theta_t^{(j)}|\Theta_t^{(i)})} & \text{if } t > 0 \end{cases} \quad (46)$$

If $i < N$ set $i = i + 1$ and go to step 3.

6. Normalize the weights. If $t < T$, set $t = t + 1$ and go to step 2.

In this notation, particles sampled from the previous population before perturbation are marked with a single asterisk (*), while perturbed particles are marked with a double asterisk (**). When $T = 1$, this reduces to the ABC rejection sampler described earlier.

In the original algorithm formulations^{51,55}, both the thresholds $\epsilon_1, \dots, \epsilon_T$ and the perturbation kernel were specified a priori. In our implementation, these quantities are adaptively chosen: the threshold at iteration t is set to the 30th percentile of the accepted distances from iteration $t - 1$, i.e., $\epsilon_t = q_{0.75}(\epsilon_{t-1})$, and the perturbation kernel for the standardized variables is Gaussian: $K_t(\Theta|\Theta^*) \sim N(\Theta_t^*, a\Sigma(\Theta_{t-1}))$ where $\Sigma(\Theta_{t-1})$ refers to the covariance matrix from the previous population and a is some scaling parameter that could either be a constant or adaptively tuned such that the the new parameter proposal Θ remains within a controlled neighborhood of the previous particle Θ^* .

The gradual convergence towards the target distribution across sample populations in SMC schemes can yield a more efficient and targeted approach for reaching the optimal distribution than rejection sampler schemes. The generation of multiple independent particles at each population also enhances robustness against becoming trapped in low-probability regions compared to MCMC algorithms. Moreover, the SMC algorithm is generally less sensitive than the rejection sampler to poorly informed prior distributions (though not completely insensitive, as it still relies on the prior for initialization and particle weight allocation in Steps 3 and 5, respectively). In addition, similar to the rejection sampler algorithm outlined earlier, the sampling and processing of multiple particle realizations within each population (Steps 2–5) can be parallelized. Furthermore, as we have demonstrated, the algorithm can be reformulated in an adaptive form, in which the thresholds $\epsilon_1, \dots, \epsilon_T$ and the perturbation kernel $K_t(\Theta|\Theta^*)$ are tuned dynamically based on the progress and behavior of the algorithm. Finally, as with the rejection sampler algorithm, additional insight into the degree of constraint for each parameter can be obtained at no extra cost through retrospective analysis of the parameter distribution evolution across populations⁵¹.

On the other hand, SMC algorithms become less practical for calibrating parameters in high-dimensional spaces due to the increased computational cost⁵¹. The reason is twofold: first, a sufficiently large number of particles must be sampled at each population iteration to adequately cover and characterize the multi-dimensional parameter space. Second, in the case of adaptive algorithms devised and used in this study, it is necessary to compute the covariance matrix and construct a multivariate perturbation kernel after each population. However, covariance matrices can become less reliable and more prone to numerical instability as the number of parameters increases. In addition, while less sensitive than the rejection sampler ABC schemes to the choice of the prior distribution, the algorithm still

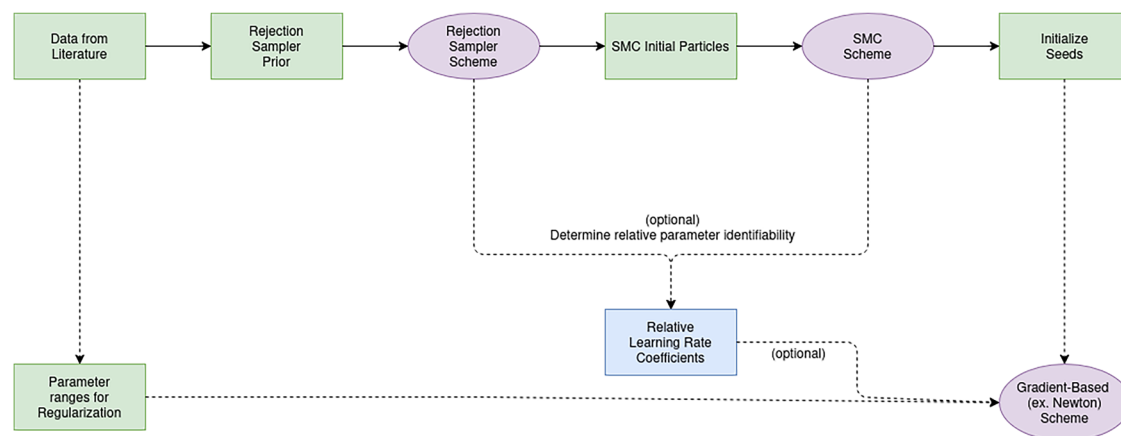


Fig. 16 | The three-stage computational pipeline is used for calibrating the ligand reaction-diffusion model from transcriptomics data. The three inference stages are shown in lilac, while the intermediate processing steps connecting them are

shown in green. Optional components (shown in blue) indicate how relative parameter identifiability inferred from the ABC stages can be used to assign relative learning rates in the gradient-based optimization step.

requires the allocation of a close and well-informed prior distribution. Finally, as in the rejection sampler algorithm, there is no direct fine-tuning of input realizations based on the output, as is the case for gradient-based optimizers.

To this end, and in the specific context of calibrating the growth factor reaction-diffusion model parameters described in “Parameter initialization and assembly,” the most effective approach is to employ the adaptive SMC algorithm described above after the allocation of an appropriate multivariate prior distribution (with a defined probability density function as required in step 5). To improve algorithm robustness and reduce computational cost, the dimensionality of the parameter space can also be decreased by fixing certain parameters a priori. Furthermore, as is the case with the rejection sampler scheme, analyzing the variances of the parameter distributions at both the initial and final populations can provide valuable insight into the degree of constraint associated with each parameter.

Parameter calibration: the computational pipeline

In “Parameter calibration schemes,” we reviewed several relevant calibration schemes and outlined the advantages and limitations of each in the context of parameter calibration for the ligand reaction-diffusion model. In this section, we combine the rejection sampler and the SMC ABC schemes with a gradient-based optimizer scheme in a sequential/hierarchical manner into a single computational pipeline, aiming to leverage the strengths of each approach while mitigating their weaknesses. The parameter calibration pipeline is illustrated in Fig. 16 and described in detail below, using the specific study case of TGF β ligand reaction-diffusion parameter calibration. All computations were carried out on a Magneto HPC cluster using multiple CPU nodes equipped with dual Intel Xeon Silver 4314 processors (32 physical cores/64 threads per node). Multi-start optimizations were parallelized across nodes and CPU cores according to the job-scheduler allocation. For the gradient-based methods, each individual seed was assigned to a single CPU core. For the ABC simulations, each Monte Carlo realization (task) was allocated three CPU cores, enabling parallel evaluation of the three TGF β isoforms in the mechanistic model.

First, the preliminary values for the parameters obtained from experimental and computational studies in the literature are used to construct a prior distribution for the rejection sampler algorithm, as described in “Parameter initialization and assembly.” This constructed prior distribution is also used to extract the permissible ranges for parameter values to be employed in constructing the regularization loss function in Eq. (39) for the gradient-based optimization algorithm. Furthermore, based on post-hoc sensitivity analysis of the parameters and their degree of constraint after the rejection sampler and SMC analyses, we can categorize the parameters into different categories based on their sensitivity and degree of constraint,

which would in turn guide the allocation of the relative learning rates for the gradient-based optimization algorithms that would benefit from the designation of the learning rates (e.g., ADAM).

Next, a population of particle realizations $N = 1000$ with the highest r are chosen from the filtered ensemble of realizations obtained from the ABC rejection sampler algorithm and used as the starting particle population for the SMC algorithm. Since the parameters span a wide range of scales and magnitudes, the SMC algorithm is conducted in the z -transformed, log parameter space; this standardization helps prevent scale dominance and numerical instability, and ensures positive values.

Next, the ABC SMC algorithm is executed with $N = 1000$ particles and $T = 40$ populations. In line with the method of sensitivity analysis by Toni et al.⁵¹ we can quantify the sensitivity of the system output to the parameters and their combinations by conducting a PCA analysis on the output of the SMC analysis, which is the ensemble of N particles from the last population. PCA provides an orthogonal transformation of the original parameter space into the new PCs, which are the eigenvectors of the corresponding covariance matrix Σ . Hence, each unit eigenvector PC_i is a linear combination of the parameters in the original space, and the corresponding eigenvalue λ_i to each eigenvector PC_i (obtained via Singular Value Decomposition) is proportional to the amount of variance $\text{var}(PC_i)$ explained by that eigenvector, or linear combination of parameters:

$$\text{var}(PC_i) = \frac{\lambda_i}{\text{trace}(\Sigma)} \quad (47)$$

In contrast to most PCA applications, where the primary focus is on PCs with the highest eigenvalues (corresponding to directions of greatest variance), our analysis focuses on the PCs with the smallest eigenvalues. These PCs represent directions in parameter space with the narrowest posterior spread, indicating parameter combinations to which the model output is most sensitive⁵¹. In other words, PCs with small eigenvalues correspond to stiff parameter combinations, while those with large eigenvalues correspond to “sloppy” parameter combinations that may not be readily identifiable^{51,57}. To explore parameter sensitivity, we first examine the composition of the two eigenvectors with the smallest eigenvalues, which capture the directions of highest output sensitivity. We then conduct a variance-reduction analysis in the z -transformed space to determine which parameters show the largest reduction in variance and compare the results of these complementary assessments.

The final stage of the parameter calibration pipeline is the gradient-based optimization. We first incorporate literature-derived parameter values to define physiologically plausible ranges, which are implemented through the regularization loss in Eq. (39). Based on the benchmarking

results comparing ADAM and Newton methods (Supplementary Note 4), we adopt a Newton-trust-region optimizer with a BFGS Hessian approximation. This choice offers second-order accuracy and avoids the need to specify learning-rate schedules.

For initialization, we draw $n = 1000$ starting points from the particles of the final SMC population. These realizations lie close to high-likelihood regions and therefore provide an effective set of seeds for probing the optimization landscape. Furthermore, in Supplementary Note 5, we describe how using starting seeds as derived from the two-stage ABC analysis improves the parameter identifiability of the MLE determined from the gradient-based Newton algorithm. Each seed is then independently optimized using the Newton-trust-region algorithm (maximum of 1000 iterations). Among the resulting optimized parameter sets, we identify the MLE as the one that achieves the minimal correlation-based loss in Eq. (40), which also minimizes the NLL (Eq. 37).

Using the computational framework to derive biological insights

We provide two illustrative examples demonstrating how the inferred parameter values can be translated into biologically meaningful insights, and how the proposed framework can guide or refine experimental hypotheses. These examples are not intended as an exhaustive biological analysis, but rather as a demonstration of the type of inference the pipeline can support.

We first compare the derived parameter values for the $TGF\beta$ production rates across the $TGF\beta$ isoforms and cell types. Before interpreting these production rates, it is important to characterize how the fibroblast subtypes differ transcriptionally. We first conduct a PCA analysis of the gene expression features and then project the FB cells and their subtype annotation onto the PCs of the highest variance, as shown in Fig. 10a, to determine whether or not the three subtypes are transcriptionally distinct. In the original transcriptomics study by Liu et al.²⁰, the authors characterized fibroblast subtypes 1, 2, and 3 as mesenchymal, pro-inflammatory, and papillary, respectively, based on their transcriptional profiles. To further validate and refine these subtype definitions within our dataset, we examined the gene loadings associated with the first three PCs and evaluated expression patterns of genes known to be associated with ECM production, chemokine/inflammatory signaling, and myofibroblast activation. Using Python's seaborn package, we calculated ECM, chemokine, and myofibroblast activity scores for each cell based on established biomarker sets, and compared these scores across the three fibroblast subtypes in Fig. 11. Next, we computed the MLEs and corresponding confidence intervals, computed under the simplifying assumption of independent homeostatic noise, for the production rates of the three $TGF\beta$ isoforms across fibroblast subtypes.

In addition, in Supplementary Note 6, we demonstrate how the framework can be used to extract mechanistic differences underlying pathological scarring, especially ones not thoroughly examined under existing analyses. We examine the effective amount of $TGF\beta$ receptors (moles/cell) between the normal day-30 wound-healing sample and the sample from AKN (subacute-chronic stage²¹).

Data availability

The scRNA and ST datasets analyzed in this study are publicly available in the GEO database under accession numbers GSE241132 and GSE206790, corresponding to the original studies by Liu et al.¹⁸ and Hong et al.²¹, respectively.

Code availability

All custom code used for data processing, model implementation, numerical solution of the reaction–diffusion system, Approximate Bayesian Computation, gradient-based optimization, and figure generation is openly available at: <https://github.com/AliDaher-the-scientist/Parameter-Calibration-pipeline-codes>. The repository contains the Python scripts used in this study. The code was written in Python 3.10 and relies on standard open-source packages, including NumPy, SciPy, PyTorch, Scanpy, Cell2location, CellPhoneDB, matplotlib, and seaborn. No restrictions apply to code access.

Received: 3 September 2025; Accepted: 22 January 2026;

Published online: 19 February 2026

References

- Hoffman, B. D., Grashoff, C. & Schwartz, M. A. Dynamic molecular processes mediate cellular mechanotransduction. *Nature* **475**, 316–323 (2011).
- Yao, L. et al. Paracrine signalling during ZEB1-mediated epithelial–mesenchymal transition augments local myofibroblast differentiation in lung fibrosis. *Cell Death Differ.* **26**, 943–957 (2019).
- Maarouf, O. H. et al. Paracrine Wnt1 drives interstitial fibrosis without inflammation by tubulointerstitial cross-talk. *J. Am. Soc. Nephrol.* **27**, 781–790 (2016).
- Park, J., Euhus, D. M. & Scherer, P. E. Paracrine and endocrine effects of adipose tissue on cancer development and progression. *Endocr. Rev.* **32**, 550–570 (2011).
- Hirano, T. IL-6 in inflammation, autoimmunity and cancer. *Int. Immunol.* **33**, 127–148 (2020).
- Shim, J. et al. Integrated analysis of single-cell and spatial transcriptomics in keloids: highlights on fibrovascular interactions in keloid pathogenesis. *J. Invest. Dermatol.* **142**, 2128–2139.e11 (2022).
- Eungdamrong, N. J. & Iyengar, R. Modeling cell signaling networks. *Biol. Cell* **96**, 355–362 (2004).
- Francis, E. A. et al. Spatial modeling algorithms for reactions and transport in biological cells. *Nat. Comput. Sci.* **5**, 76–89 (2025).
- Rangamani, P. & Iyengar, R. Modelling cellular signalling systems. *Essays Biochem.* **45**, 83–94 (2008).
- Murphy, K. E., Hall, C. L., Maini, P. K., McCue, S. W. & McElwain, D. L. S. A fibrocontractive mechanochemical model of dermal wound closure incorporating realistic growth factor kinetics. *Bull. Math. Biol.* **74**, 1143–1170 (2012).
- Turing, A. The chemical basis of morphogenesis. *Philos. Trans. R. Soc. Lond. Ser. B Biol. Sci.* **237**, 37–72 (1952).
- Pasetto, S. et al. Calibrating tumor growth and invasion parameters with spectral spatial analysis of cancer biopsy tissues. *npj Syst. Biol. Appl.* **10**, 1–9 (2024).
- Madamanchi, A., Ingle, M., Hinck, A. P. & Umlis, D. M. Computational modeling of TGF- β 2:TBRI:TBRII receptor complex assembly as mediated by the TGF- β coreceptor betaglycan. *Biophys. J.* **122**, 1342–1354 (2023).
- De Crescenzo, G., Grothe, S., Zwaagstra, J., Tsang, M. & O'Connor-McCourt, M. D. Real-time monitoring of the interactions of transforming growth factor- β (TGF- β) isoforms with latency-associated protein and the ectodomains of the TGF- β type II and III receptors reveals different kinetic models and stoichiometries of binding. *J. Biol. Chem.* **276**, 29632–29643 (2001).
- De Crescenzo, G., Pham, P. L., Durocher, Y. & O'Connor-McCourt, M. D. Transforming growth factor- β (TGF- β) binding to the extracellular domain of the type II TGF- β receptor: receptor capture on a biosensor surface using a new coiled-coil capture system demonstrates that avidity contributes significantly to high affinity binding. *J. Mol. Biol.* **328**, 1173–1183 (2003).
- Radaev, S. et al. Ternary complex of transforming growth factor- β 1 reveals isoform-specific ligand recognition and receptor recruitment in the superfamily. *J. Biol. Chem.* **285**, 14806–14814 (2010).
- Huang, T., Schor, S. L. & Hinck, A. P. Biological activity differences between TGF- β 1 and TGF- β 3 correlate with differences in the rigidity and arrangement of their component monomers. *Biochemistry* **53**, 5737–5749 (2014).
- Liu, Z. et al. Spatiotemporal single-cell roadmap of human skin wound healing. *Cell Stem Cell* **32**, 479–498.e8 (2025).
- Massagué, J., Cheifetz, S., Boyd, F. T. & Andres, J. L. TGF- β Receptors and TGF- β Binding Proteoglycans: Recent Progress in

- Identifying Their Functional Properties. *Ann. N. Y. Acad. Sci.* **593**, 59–72 (1990).
20. Liu, Z., Li, D. & Landén, N. X. 538 spatiotemporal single-cell roadmap of human skin wound healing. *J. Investig. Dermatol.* **144**, S321 (2024).
21. Hong, Y.-K. et al. Profibrotic subsets of SPP1+ macrophages and POSTN+ fibroblasts contribute to fibrotic scarring in acne keloidalis. *J. Investig. Dermatol.* **144**, 1491–1504.e10 (2024).
22. Kim, K. K., Sheppard, D. & Chapman, H. A. TGF- β 1 signaling and tissue fibrosis. *Cold Spring Harb. Perspect. Biol.* **10**, a022293 (2018).
23. Raue, A. et al. Structural and practical identifiability analysis of partially observed dynamical models by exploiting the profile likelihood. *Bioinformatics* **25**, 1923–1929 (2009).
24. Hass, H. et al. Benchmark problems for dynamic modeling of intracellular processes. *Bioinformatics* **35**, 3073–3082 (2019).
25. Villaverde, A. F., Fröhlich, F., Weindl, D., Hasenauer, J. & Banga, J. R. Benchmarking optimization methods for parameter estimation in large kinetic models. *Bioinformatics* **35**, 830–838 (2019).
26. Ramilowski, J. A. et al. A draft network of ligand–receptor-mediated multicellular signalling in human. *Nat. Commun.* **6**, 7866 (2015).
27. Cang, Z. et al. Screening cell–cell communication in spatial transcriptomics via collective optimal transport. *Nat. Methods* **20**, 218–228 (2023).
28. Gleizes, P. E. et al. TGF- β latency: biological significance and mechanisms of activation. *Stem Cells* **15**, 190–197 (1997).
29. Daher, A. & Payne, S. A network-based model of dynamic cerebral autoregulation. *Microvasc. Res.* **147**, 104503 (2023).
30. Daher, A. Using approximate Bayesian computation to calibrate the model parameters characterizing the autoregulatory behavior of microvessels. *Int. J. Numer. Methods Biomed. Eng.* **41**, e70023 (2025).
31. Efremova, M., Vento-Tormo, M., Teichmann, S. A. & Vento-Tormo, R. CellPhoneDB: inferring cell–cell communication from combined expression of multi-subunit ligand–receptor complexes. *Nat. Protoc.* **15**, 1484–1506 (2020).
32. Vento-Tormo, R. et al. Single-cell reconstruction of the early maternal–fetal interface in humans. *Nature* **563**, 347–353 (2018).
33. Jin, S., Plikus, M. V. & Nie, Q. CellChat for systematic analysis of cell–cell communication from single-cell transcriptomics. *Nat. Protoc.* **20**, 180–219 (2025).
34. Liu, Z., Sun, D. & Wang, C. Evaluation of cell–cell interaction methods by integrating single-cell RNA sequencing data with spatial information. *Genome Biol.* **23**, 218 (2022).
35. Luo, J., Deng, M., Zhang, X. & Sun, X. ESICCC as a systematic computational framework for evaluation, selection, and integration of cell–cell communication inference methods. *Genome Res.* **33**, 1788–1805 (2023).
36. Hiura, T., Seno, S. & Matsuda, H. Inference of cell–cell interactions through spatial transcriptomics data using graph convolutional neural networks. in *Parallel and Distributed Processing Techniques* (eds Arabnia, H. R. et al.) 139–152 (Springer Nature Switzerland, 2025).
37. Zhao, M. et al. Targeting fibrosis: mechanisms and clinical trials. *Signal Transduct. Target. Ther.* **7**, 206 (2022).
38. Cromack, D. T. et al. Transforming growth factor beta levels in rat wound chambers. *J. Surg. Res.* **42**, 622–628 (1987).
39. Yang, L., Qiu, C. X., Ludlow, A., Ferguson, M. W. J. & Brunner, G. Active transforming growth factor- β 1 wound repair. *Am. J. Pathol.* **154**, 105–111 (1999).
40. Shan, M. & Wang, Y. Viewing keloids within the immune microenvironment. *Am. J. Transl. Res.* **14**, 718–727 (2022).
41. Wolock, S. L., Lopez, R. & Klein, A. M. Scrublet: computational identification of cell doublets in single-cell transcriptomic data. *Cell Syst.* **8**, 281–291.e9 (2019).
42. Wolf, F. A., Angerer, P. & Theis, F. J. SCANPY: large-scale single-cell gene expression data analysis. *Genome Biol.* **19**, 15 (2018).
43. Ji, A. L. et al. Multimodal analysis of composition and spatial architecture in human squamous cell carcinoma. *Cell* **182**, 1661–1662 (2020).
44. Visium Spatial Platform. <https://www.10xgenomics.com/platforms/visium>.
45. Kleshchevnikov, V. et al. Cell2location maps fine-grained cell types in spatial transcriptomics. *Nat. Biotechnol.* **40**, 661–671 (2022).
46. Welcome to CellPhoneDB's documentation. - cellphonedb documentation. <https://cellphonedb.readthedocs.io/en/latest/index.html>.
47. Dell'Acqua, G. & Bersani, A. M. Quasi-steady state approximations and multistability in the double phosphorylation–dephosphorylation cycle. in *Biomedical Engineering Systems and Technologies* (eds Fred, A., Filipe, J. & Gamboa, H.) 155–172 (Springer, 2013).
48. Schmiester, L., Schälte, Y., Fröhlich, F., Hasenauer, J. & Weindl, D. Efficient parameterization of large-scale dynamic models based on relative measurements. *Bioinformatics* **36**, 594–602 (2019).
49. Deng, Z. et al. TGF- β signaling in health, disease and therapeutics. *Signal Transduct. Target. Ther.* **9**, 1–40 (2024).
50. PyTorch. <https://pytorch.org/>.
51. Toni, T., Welch, D., Strelkowa, N., Ipsen, A. & Stumpf, M. P. Approximate Bayesian computation scheme for parameter inference and model selection in dynamical systems. *J. R. Soc. Interface* **6**, 187–202 (2009).
52. Pritchard, J. K., Seielstad, M. T., Perez-Lezaun, A. & Feldman, M. W. Population growth of human Y chromosomes: a study of Y chromosome microsatellites. *Mol. Biol. Evolution* **16**, 1791–1798 (1999).
53. Daher, A. & Payne, S. A dynamic multiscale model of cerebral blood flow and autoregulation in the microvasculature. *Appl. Math. Model.* **123**, 213–240 (2023).
54. Marjoram, P., Molitor, J., Plagnol, V. & Tavaré, S. Markov chain Monte Carlo without likelihoods. *Proc. Natl. Acad. Sci. USA* **100**, 15324–15328 (2003).
55. Sisson, S. A., Fan, Y. & Tanaka, M. M. Sequential Monte Carlo without likelihoods. *Proc. Natl. Acad. Sci. USA* **104**, 1760–1765 (2007).
56. Daher, A. & Payne, S. The conducted vascular response as a mediator of hypercapnic cerebrovascular reactivity: a modelling study. *Comput. Biol. Med.* **170**, 107985 (2024).
57. Gutenkunst, R. N. et al. Universally sloppy parameter sensitivities in systems biology models. *PLOS Comput. Biol.* **3**, e189 (2007).
58. Diffusion of molecules and cytokines. <https://personal.math.ubc.ca/ais/website/status/diffuse.html>.
59. Maroudas, A. Distribution and diffusion of solutes in articular cartilage. *Biophys. J.* **10**, 365–379 (1970).
60. Matharu, Z. et al. Detecting transforming growth factor- β release from liver cells using an aptasensor integrated with microfluidics. *Anal. Chem.* **86**, 8865–8872 (2014).
61. Sun, T., Adra, S., Smallwood, R., Holcombe, M. & MacNeil, S. Exploring hypotheses of the actions of TGF- β 1 in epidermal wound healing using a 3D computational multiscale model of the human epidermis. *PLoS ONE* **4**, e8515 (2009).
62. Fuchs, E. Epidermal differentiation: the bare essentials. *J. Cell Biol.* **111**, 2807–2814 (1990).
63. Hinck, A. P., Mueller, T. D. & Springer, T. A. Structural biology and evolution of the TGF- β family. *Cold Spring Harb. Perspect. Biol.* **8**, a022103 (2016).
64. Baardsnes, J., Hinck, C. S., Hinck, A. P. & O'Connor-McCourt, M. D. TBR-II discriminates the high- and low-affinity TGF- β isoforms via two hydrogen-bonded ion pairs. *Biochemistry* **48**, 2146–2155 (2009).
65. Needleman, B. W., Choi, J., Burrows-Mezu, A. & Fontana, J. A. Secretion and binding of transforming growth factor beta by scleroderma and normal dermal fibroblasts. *Arthritis Rheum.* **33**, 650–656 (1990).

66. Brandes, M. E., Wakefield, L. M. & Wahl, S. M. Modulation of monocyte type I transforming growth factor-beta receptors by inflammatory stimuli. *J. Biol. Chem.* **266**, 19697–19703 (1991).
67. Wahl, S. M. et al. Transforming growth factor type beta induces monocyte chemotaxis and growth factor production. *Proc. Natl. Acad. Sci. USA* **84**, 5788–5792 (1987).
68. Fafeur, V., Terman, B. I., Blum, J. & Böhlen, P. Basic FGF treatment of endothelial cells down-regulates the 85-KDa TGF beta receptor subtype and decreases the growth inhibitory response to TGF-beta 1. *Growth Factors* **3**, 237–245 (1990).

Acknowledgements

The authors would like to thank Nicholas Dawes, Scientific Computing Engineer at the University of Dundee, for his guidance and support with submitting the computational runs on the Dundee HPC. This work was supported by the Agence Nationale de la Recherche (ANR) under grant number ANR-21-CE45-0025-01.

Author contributions

A.D. conceived and designed the study, developed the computational code, performed the analyses, drafted and revised the manuscript, and conducted the benchmarking test cases for the revision. R.E. and D.T. provided supervision, guidance, and feedback throughout the study. R.E. secured funding, and D.T. provided computational resources. All authors reviewed and approved the final manuscript.

Competing interests

The authors declare no competing interests.

Additional information

Supplementary information The online version contains supplementary material available at <https://doi.org/10.1038/s41540-026-00657-8>.

Correspondence and requests for materials should be addressed to Ali Daher.

Reprints and permissions information is available at <http://www.nature.com/reprints>

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Open Access This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

© The Author(s) 2026, modified publication 2026