
KINDLE: Knowledge-Guided Distillation for Prior-Free Gene Regulatory Network Inference

Rui Peng^{1,2}, Yuchen Lu³, Qichen Sun¹, Yuxing Lu¹, Chi Zhang¹, Ziru Liu⁴, Jinzhuo Wang^{1*}

¹Department of Big Data and Biomedical AI, College of Future Technology, Peking University

²Center for BioMed-X Research, Academy for Advanced Interdisciplinary Studies, Peking University

³School of Physics, Peking University

⁴Yuanpei College, Peking University

{pengrui, luyuchen2002, 2000010820, luyx, cszc21, lzh}@stu.pku.edu.cn

wangjinzhuo@pku.edu.cn

Abstract

Gene regulatory network (GRN) inference serves as a cornerstone for deciphering cellular decision-making processes. Early approaches rely exclusively on gene expression data, thus their predictive power remain fundamentally constrained by the vast combinatorial space of potential gene-gene interactions. Subsequent methods integrate prior knowledge to mitigate this challenge by restricting the solution space to biologically plausible interactions. However, we argue that the effectiveness of these approaches is contingent upon the precision of prior information and the reduction in the search space will circumscribe the models' potential for novel biological discoveries. To address these limitations, we introduce KINDLE, a three-stage framework that decouples GRN inference from prior knowledge dependencies. KINDLE trains a teacher model that integrates prior knowledge with temporal gene expression dynamics and subsequently distills this encoded knowledge to a student model, enabling accurate GRN inference solely from expression data without access to any prior. KINDLE achieves state-of-the-art performance across four benchmark datasets. Notably, it successfully identifies key transcription factors governing mouse embryonic development and precisely characterizes their functional roles. In mouse hematopoietic stem cell data, KINDLE accurately predicts fate transition outcomes following knockout of two critical regulators (Gata1 and Spi1). These biological validations demonstrate our framework's dual capability in maintaining topological inference precision while preserving discovery potential for novel biological mechanisms.

1 Introduction

Gene regulatory network (GRN) represents a directed graph that depicts the regulatory interactions between genes, where nodes consist of transcription factors (TFs) and target genes (TGs). A directed edge between a TF and a TG signifies the TF's capacity to bind the cis-regulatory elements of the TG and subsequently modulates its transcriptional activity [1]. GRN provides mechanistic blueprints for understanding regulatory logic underlying lineage commitment, maintenance, and reprogramming [2]. Precisely resolved GRN enables mechanistic interpretations of lineage bifurcation, aging process, and tumor-related dysregulation [3].

Despite their biological significance, GRN inference remains technically challenging. Early inference methods that rely solely on gene expression data face inherent limitations: The explorable TF-TG

*Corresponding Author

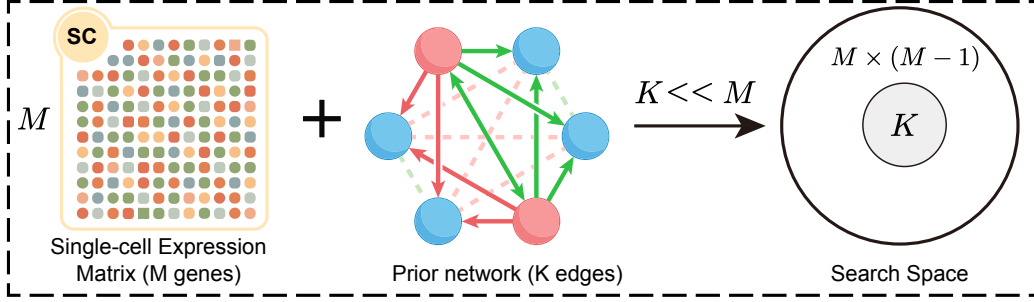


Figure 1: For a dataset with M genes, the search space for gene pairs spans $M \times (M - 1)$ possible interactions, as self-loop edges are not considered. Prior-based methods drastically narrow the exploration to the K edges supported by prior knowledge ($K \ll M$).

interaction space scales quadratically with the number of genes, resulting in approximately 1 billion potential regulatory interactions within the whole genome (comprising approximate 30,000 genes). This vast search space fundamentally constrains the performance of expression-based approaches. Contemporary methods address this by incorporating prior knowledge from complementary data (e.g., scATAC-seq [4] or Hi-C [5]) to constrain the search space to pre-defined TF-TG pairs, as illustrated in Figure 1. Although prior-based approaches enhance performance by narrowing the search space, they impose two major limitations. Firstly, with a fixed prior network, an algorithm is confined to searching among its existing edges, and its performance depends on the overlap between the prior and the ground truth network. A perfect match allows for 100% accuracy, while minimal or no overlap leads to zero accuracy for all algorithms. Secondly, limiting candidate edges to the prior prevents the detection of regulatory interactions absent from it, a critical drawback that fundamentally constrains a model’s utility for scientific discovery. For example, analyzing gene expression data from cancer cells might reveal a previously unknown transcription factor regulating a pro-oncogenic gene network. Such a discovery, which is impossible when confined to a prior network of already-validated interactions, could lead to the development of new precision therapies.

To overcome prior-dependent limitations, we propose a strategy inspired by learning with privileged information [6]. This paradigm obtains a teacher model using supplementary privileged features during training, followed by knowledge transfer to a student model operating without access to such features. Building on this framework, we develop a three-stage architecture named KINDLE (**K**nowledge-gu**I**ded Network **D**isti**L**lation for prior-free GRN inf**E**rence) to infer accurate GRN without relying on prior information. The first stage trains a teacher model integrating both gene expression data and external priors. Notably, inspired by TRIGON’s temporal causality modeling [7], the teacher model explicitly captures temporal regulatory dynamics by predicting future gene expression states from historical expression profiles, rather than relying on static gene co-expression analysis. The incorporated prior knowledge further refines the candidate regulatory space, generating temporally coherent and biologically plausible regulatory maps. The second stage implements knowledge distillation to train a student model through teacher supervision while completely eschewing prior information. The final stage deploys the student model for prior-independent GRN inference using expression data exclusively, thereby achieving scalable and unbiased reconstruction of regulatory networks. Our contributions are summarized as follows:

- We propose KINDLE to eliminate prior dependence in GRN inference by knowledge distillation, which achieves state-of-the-art performance across four benchmark datasets without requiring prior knowledge.
- On mouse embryonic stem cell development data, KINDLE successfully identifies key TFs and predicts their functional roles during differentiation processes.
- For mouse hematopoietic stem cell development, KINDLE accurately predicts the effects of Gata1 and Spi1 knockouts on cell fate determination, demonstrating its capability to capture critical regulatory mechanisms.

2 Related work

Prior-based GRN inference. Early attempts in GRN inference predominantly relied on co-expression analyses from bulk or single-cell transcriptomic datasets [8–13]. However, the inherent limitations of unimodal data approaches became evident due to the vast combinatorial search space of potential TF–TG interactions, which severely restricted their predictive performance. To constrain the solution space, contemporary computational pipelines strategically integrated external biological priors during model optimization. For instance, LINGER [14] employed a neural network architecture trained on paired single-cell RNA-seq and ATAC-seq profiles to predict gene expression dynamics through systematic integration of TF abundance and chromatin accessibility. CEFCON [15] implemented a graph attention network initialized with motif-informed adjacency matrix, synergistically coupling cell lineage-specific GRN inference with network control theory. The Celloracle framework [16] operationalized promoter-enhancer interaction maps coupled with DNA motif annotations to establish a base GRN architecture, which undergoes iterative refinement through ridge regression. While demonstrating methodological innovation, all these approaches exhibited fundamental dependence on the precision and comprehensiveness of incorporated prior knowledge, inaccuracies in prior specification risk propagating systematic biases.

Privileged-feature distillation. Privileged-feature distillation uses auxiliary data accessible exclusively during training while eliminating their requirement during downstream deployment. In this paradigm, a teacher model with privileged input creates informative soft targets or latent representations to supervise a student model restricted to privileged input. Theoretical analyses in learning-to-rank contexts showed that balancing data-driven loss and teacher guidance helps distilled students outperform non-distilled models [6]. Empirically, the BLEND computational framework [17] validated this by applying the methodology to large-scale neurobiological datasets, where behavioral trajectory data acted as privileged supervisory signals during teacher model optimization, and the distilled neural-activity-only student excelled in population coding decryption tasks. Overall, these theoretical and applied advancements established privileged-feature distillation as a robust way to eliminate the dependence on noisy or resource-intensive prior information, a critical but underexploited property with potential to enhance GRN inference.

3 Methodology

3.1 Theoretical Foundation

Our framework is grounded on the causal hypothesis that GRN intrinsically govern transcriptional dynamics through time-evolving interactions. Formally, let $\mathbf{G} \in \mathbb{R}^{N \times M}$ denotes the temporal single-cell expression matrix, where N represents temporally ordered cellular states and M is the number of genes. We posit that an accurate GRN adjacency matrix $\mathbf{A} \in \mathbb{R}^{M \times M}$ should encode sufficient mechanistic information to predict future expression states from historical observations and thus satisfy:

$$\mathbf{G}_{T+1:T+W} \approx \mathcal{F}(\mathbf{G}_{1:T}, \mathbf{A}) \quad (1)$$

where $\mathcal{F}$ encodes the nonlinear regulatory kinetics, T and W define historical and future time windows respectively. The GRN inference problem is thus reframed as learning a minimal sufficient interaction matrix $\mathbf{A}^*$ that minimizes the difference between predicted and actual gene expression profiles:

$$\mathbf{A}^* = \arg \min_{\mathbf{A}} \|\mathcal{F}(\mathbf{G}_{1:T}, \mathbf{A}) - \mathbf{G}_{T+1:T+W}\|_2^2 \quad (2)$$

We present KINDLE to operationalize this theory and infer the prior-free GRN by three sequential phases: initial supervised training of the teacher model to assimilate prior network guidance, subsequent distillation of the teacher’s regulatory insight into a lightweight student model, and ultimately deployment of the prior independent student model for high-fidelity GRN inference.

3.2 Teacher Model

As illustrated in Figure 2, the teacher model is equipped with hierarchical attention mechanisms, consisting of temporal and spatial layers. In the temporal layer, a lower triangular mask is applied to the attention weights, ensuring each gene’s expression at time step t only attends to its historical

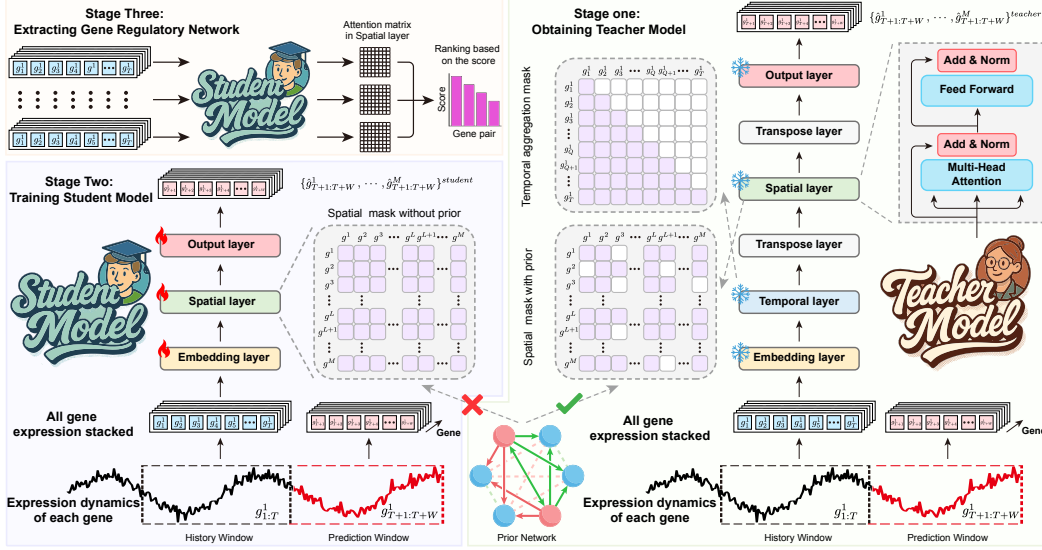


Figure 2: Illustration of KINDLE framework. The pipeline consists of three stages: (1) Teacher model integrates prior knowledge to learn causal regulatory relationships via explicit modeling of expression state transitions across time windows. (2) With teacher model parameters frozen, knowledge distillation transfers regulatory insights to a lightweight student model that operates exclusively on expression data, free of prior inputs. (3) The trained student model is deployed to infer GRN, yielding prior-decoupled network that maintain high accuracy.

states $\{1, \dots, t-1\}$. This causal constraint mirrors the irreversibility of cellular differentiation, where progenitor cells cannot access transcriptional information from their descendants. The spatial layer employs a prior-derived binary mask $\mathbf{M}^{spatial} \in \{0, 1\}^{M \times M}$, where $\mathbf{M}_{ij}^{spatial} = 1$ indicates a documented regulatory interaction from gene i to j in prior knowledge. This mask sparsifies attention computation by restricting cross-gene interactions to curated regulatory pairs, effectively pruning unvalidated relationships while preserving interpretability. Architecturally, the temporal layer processes input tensors $\mathbf{X} \in \mathbb{R}^{B \times T \times M}$ (batch size B , time steps T , genes M) and outputs a tensor of identical dimensions. To enable gene-centric regulatory modeling in the subsequent spatial layer, we perform axis transposition $\mathbb{R}^{B \times T \times M} \rightarrow \mathbb{R}^{B \times M \times T}$, restructuring the tensor to treat each gene's temporal trajectory as an independent sequence. This dimensional reorganization permits parallelized computation of gene-specific attention weights across all M genes while maintaining temporal dependencies. Following spatial attention computation, the tensor undergoes transposition operation $\mathbb{R}^{B \times M \times T} \rightarrow \mathbb{R}^{B \times T \times M}$, followed by linear projection to $\mathbb{R}^{B \times W \times M}$ for W -step gene expression prediction. The end-to-end framework optimizes regulatory dynamics by minimizing the mean squared error:

$$\mathcal{L}_{\text{teacher}} = \frac{1}{B \cdot W \cdot M} \sum_{b=1}^B \sum_{w=1}^W \sum_{m=1}^M \left\| \hat{\mathbf{Y}}_{t+w,m}^{(b)} - \mathbf{Y}_{t+w,m}^{(b)} \right\|_2^2 \quad (3)$$

where $\hat{\mathbf{Y}}$ and $\mathbf{Y}$ denote predicted and ground truth expression matrix respectively, indexed by batch b , forecast window w , and gene m . While the attention matrix $\mathbf{A} \in \mathbb{R}^{M \times M}$ extracted from teacher model's spatial layer during inference could be a prior-constrained approximation of the theoretically optimal matrix $\mathbf{A}^*$ defined in Eq.2, the solution remains fundamentally constrained by its prior-dependent architecture. Specifically, the binary masking operation irreversibly eliminates attention weights for gene pairs absent in the prior knowledge (i.e., positions where $\mathbf{M}_{ij}^{spatial} = 0$), thereby restricting the teacher model's attention exclusively to a sparse subset of regulatory interactions defined by prior-informed positions (i.e., $\mathbf{M}_{ij}^{spatial} = 1$). This prior-induced myopia severely limits applicability to emerging biological systems with incomplete interactome annotations. To overcome this fundamental limitation, we design a student model that learns the teacher's regulatory knowledge through distillation without inheriting its prior constraints.

3.3 Student Model

We formalize the student model as $f_{\theta_S} \in \{f | f : \mathbf{G}_{1:T} \mapsto \mathbf{G}_{T+1:T+W}\}$, operating exclusively on raw expression matrix $\mathbf{G}_{1:T} \in \mathbb{R}^{B \times T \times M}$ without prior network integration. We let f_{θ_T} be the teacher model and the parameter optimization of the student model aims to minimize the following loss:

$$\underbrace{\alpha \cdot \sum_{b,w,m} \|f_{\theta_S}(\mathbf{G}_{1:T})_{b,m} - \mathbf{G}_{b,T+w,m}\|_2^2}_{\text{Prediction Loss}} + \underbrace{(1 - \alpha) \cdot \sum_{b,m} \mathcal{L}_{\text{distill}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m})}_{\text{Regulatory Distillation Loss}} \quad (4)$$

The hyperparameter $\alpha \in (0, 1)$ governs the trade-off between expression prediction fidelity and regulatory knowledge transfer. Crucially, as diagrammed in Figure 2, the student architecture implements two critical modifications: (1) Elimination of the prior-dependent spatial mask $\mathbf{M}^{\text{spatial}}$ in attention computation, enabling unrestricted interaction modeling between all gene pairs. (2) Removal of the teacher’s temporal layer while preserving temporal causality through distillation, resulting in a lightweight model.

Distinct formulations of the distillation loss $\mathcal{L}_{\text{distill}}$ can extract different dimensions of the teacher model’s knowledge. In the course of this research, we explore four primary distillation strategies within our KINDLE framework. Each of these strategies is meticulously designed to convey specific aspects of the teacher model’s knowledge to the student model, thereby enhancing the latter’s performance and understanding.

Hard Distillation. We optimize this baseline through a predictive congruence objective, where $\mathcal{L}_{\text{distill}}$ is constructed as direct predictive alignment. Specifically, the framework achieves this by optimizing the squared L2-norm divergence between the teacher’s terminal predictions and the student’s corresponding outputs, enforcing knowledge transfusion via deterministic supervision of final-layer activations:

$$\mathcal{L}_{\text{distill}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) = \|f_{\theta_S}(\mathbf{G}_{1:T})_{b,m} - f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}\|_2^2 \quad (5)$$

Soft Distillation. This paradigm implements probabilistic knowledge transfer through entropy-regulated distribution matching. The framework introduces temperature parameter τ to soften the logits before applying the softmax function, formally expressed as:

$$\mathcal{L}_{\text{distill}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) = \text{KL} \left(\sigma \left(\frac{f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}}{\tau} \right) \parallel \sigma \left(\frac{f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}}{\tau} \right) \right) \quad (6)$$

where σ is the softmax function, and $\text{KL}(\cdot \parallel \cdot)$ is the Kullback-Leibler divergence.

In addition to the aforementioned hard target distillation and soft probabilistic matching, we develop correlation distillation to preserve structural dependencies in feature representations. The core objective is to align the teacher-student correlation manifolds through kernel-induced similarity measures, formalized as:

$$\mathcal{L}_{\text{distill}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) = \mathcal{K}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) \quad (7)$$

where $\mathcal{K}(\cdot, \cdot)$ denote kernel methods to compute the correlation between output of f_{θ_S} and f_{θ_T} . To address the challenges posed by the high dimensionality of embedded feature spaces in analyzing complex inter-instance correlations, we propose two different kernel methods to effectively capture the high-order correlations between instances within the feature space.

Bilinear Pool. It computes inter-instance correlations through outer product operations. Formally, the Bilinear Pool kernel is defined as:

$$\mathcal{K}_{\text{bilinear}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) = (f_{\theta_S}(\mathbf{G}_{1:T})_{b,m})^\top (f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) \quad (8)$$

Gaussian RBF. This non-linear operator characterizes instance relationships through exponentially decaying similarity metrics, possessing stronger non-linear manifold learning capabilities compared

Table 1: Comparison of the proposed KINDLE framework with other GRN inference methods on four datasets provided by BEELINE [18]. **Bold values** denote the best performance for the corresponding metric.

Methods		mESC			mHSC-E			mHSC-L			mHSC-GM		
		AUROC	AUPRC	F1	AUROC	AUPRC	F1	AUROC	AUPRC	F1	AUROC	AUPRC	F1
Without Prior	GRNBoost2 [8]	0.537	0.127	0.203	0.397	0.034	0.087	0.515	0.181	0.297	0.474	0.083	0.146
	GENIE3 [9]	0.545	0.137	0.218	0.381	0.042	0.108	0.486	0.183	0.322	0.437	0.078	0.162
	Random	0.506	0.083	0.152	0.493	0.087	0.161	0.518	0.135	0.227	0.504	0.083	0.154
Prior Guided	CEFCON [15]	0.479	0.253	0.429	0.531	0.405	0.551	0.653	0.659	0.675	0.457	0.444	0.647
	Celloracle [16]	0.490	0.177	0.305	0.536	0.290	0.420	0.557	0.277	0.368	0.487	0.243	0.401
	NetREX [19]	0.522	0.128	0.217	0.511	0.117	0.211	0.520	0.177	0.282	0.526	0.144	0.219
	Prior_Random	0.498	0.318	0.482	0.492	0.389	0.570	0.522	0.551	0.691	0.509	0.464	0.627
KINDLE	KINDLE (Soft distillation)	0.747	0.636	0.519	0.561	0.559	0.691	0.599	0.670	0.752	0.562	0.789	0.864
	KINDLE (Hard distillation)	0.753	0.643	0.526	0.564	0.578	0.711	0.599	0.669	0.757	0.569	0.793	0.871
	KINDLE (Bilinear Pool)	0.751	0.644	0.521	0.551	0.574	0.723	0.567	0.581	0.761	0.561	0.787	0.867
	KINDLE (Gaussian RBF)	0.757	0.646	0.529	0.594	0.601	0.731	0.600	0.672	0.763	0.570	0.799	0.875

to bilinear methods. The kernel admits low-rank Taylor series approximation while preserving topological structures in feature space. Formally, the Gaussian RBF kernel is defined as:

$$\mathcal{K}_{\text{gaussian}}(f_{\theta_S}(\mathbf{G}_{1:T})_{b,m}, f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}) = \exp\left(-\frac{\|f_{\theta_S}(\mathbf{G}_{1:T})_{b,m} - f_{\theta_T}(\mathbf{G}_{1:T}, \mathbf{M}^{\text{spatial}})_{b,m}\|_2^2}{2\lambda^2}\right) \quad (9)$$

where λ is a hyperparameter that controls the width of the gaussian function.

3.3.1 GRN Inference

Given the input gene expression time series $\mathbf{G} \in \mathbb{R}^{N \times M}$, where N denotes the temporal sequence length and M represents the number of genes, we partition the sequence into segments of length $T \in \mathbb{N}^+$. Under the divisibility condition $T \mid N$, we obtain $H = \frac{N}{T}$ non-overlapping samples $\{\mathcal{S}^{(g)}\}_{g=1}^H$, each containing T consecutive temporal observations:

$$\mathcal{S}^{(g)} = \mathbf{G}_{(T \cdot (g-1) + 1):(T \cdot g)} \quad \forall g \in \{1, \dots, H\} \quad (10)$$

For each partitioned sample $\mathcal{S}^{(g)}$, the student model f_{θ_S} generates attention matrix $\mathbf{A}^{(g)} \in \mathbb{R}^{M \times M}$ in its spatial layer. We compute the optimal approximation $\hat{\mathbf{A}}$ to the theoretical $\mathbf{A}^*$ in Eq.2 through temporal ensemble:

$$\hat{\mathbf{A}} = \frac{1}{H} \sum_{g=1}^H \mathbf{A}^{(g)} \quad (11)$$

The final GRN $\mathcal{G}_{pred}$ is established by ranking the edge weights in the matrix $\hat{\mathbf{A}}$ and selecting the top- k most significant connections, where k corresponds exactly to the number of edges in the ground truth regulatory network $\mathcal{G}_{gt}$ provided with each benchmarking dataset:

$$\mathcal{G}_{pred} = \{(i, j) \mid \hat{A}_{ij} \in \text{Top}_k(\hat{\mathbf{A}})\}, \quad k = |\mathcal{G}_{gt}| \quad (12)$$

The detailed pseudocode implementations of KINDLE are provided in Appendix F.

4 Experiments

4.1 KINDLE Achieved State-of-the-Art Performance in GRN Benchmarks

The evaluation of KINDLE strictly adheres to the benchmarking protocol introduced in BEELINE [18]. We systematically validated our approach on four mouse differentiation datasets, the embryonic stem cell (mESC) dataset as well as three hematopoietic lineages: Erythrocyte (mHSC-E), Granulocyte-Monocyte (mHSC-GM), and Lymphocyte (mHSC-L). Following BEELINE’s established framework, we treated GRN inference as a binary classification task, employing lineage-matched reference GRN derived from ChIP-seq experiments as ground truth (see Appendix B.1 for detailed information). Performance was quantified through three metrics: area under the receiver operating characteristic curve (AUROC), area under the precision-recall curve (AUPRC), and F1 score (see Appendix B.2 for pseudocode of metric calculation). KINDLE was compared against seven competitive baselines: expression-based methods (GENIE3 [9], GRNBoost2 [8]), prior-based

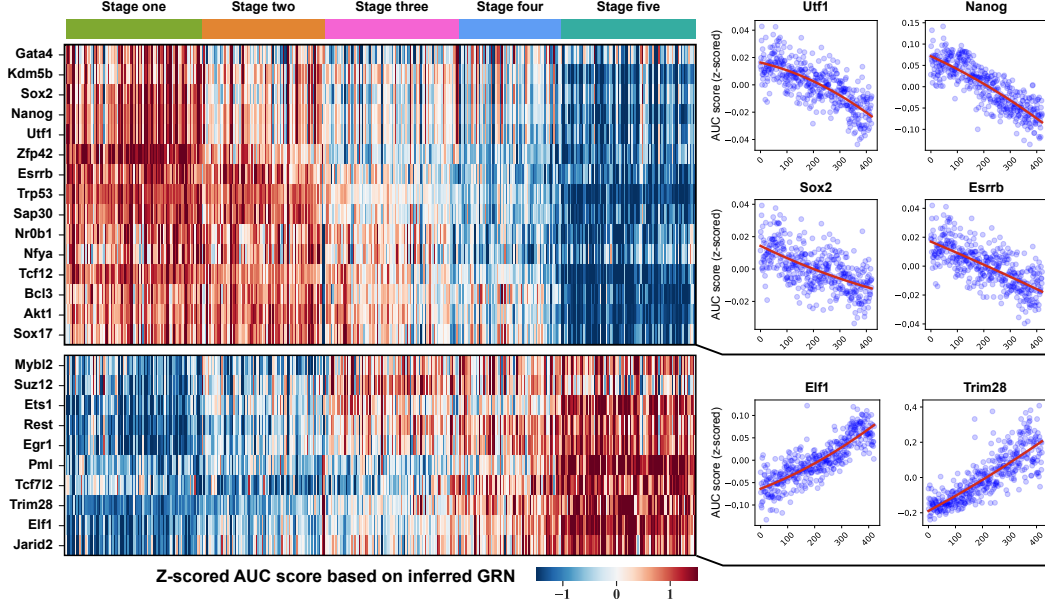


Figure 3: Temporal dynamics of TF regulatory performance during mouse embryonic stem cell differentiation. **Left:** Heatmap visualization of z-scored AUC scores reveals bimodal temporal patterns through hierarchical clustering. Two distinct TF clusters emerge, demonstrating their divergent regulatory roles during differentiation. **Right:** Trend analysis of AUC score, with developmental time points on x-axis and z-scored AUC score on y-axis. Red curves represent quadratic polynomial fits to the dynamic profiles. Complete regression curves for all 25 TFs are available in Appendix E.

approaches (CEFCON [15], CellOracle [16], NetREX [19]), and two random controls (Random, Prior_Random). Detailed descriptions of datasets, ground truth networks, and baseline models are provided in Appendix A.1, Appendix A.2, and Appendix A.3 respectively.

As shown in Table 1, all four KINDLE variants demonstrated substantial improvements over baselines despite requiring no external biological priors. The soft distillation variant, representing our weakest configuration, surpassed the best expression-based method (GENIE3) by 0.499, 0.517, 0.490, 0.711 in AUPRC across datasets and outperformed CellOracle by 0.546 in mHSC-GM. Among variants, the one with Gaussian RBF (hereafter **KINDLE-Gaussian**) achieved the best performance in 11 of 12 dataset-metric combinations. On the mESC dataset, KINDLE-Gaussian improved AUROC from 0.545 (GENIE3) to 0.757 (39% increase), elevated AUPRC from 0.253 (CEFCON) to 0.646 (155% improvement), and raised F1 score from 0.429 to 0.529 (23% gain). Comparable enhancements emerged in hematopoietic lineages: AUPRC increased by 48% (0.405 $\rightarrow$ 0.601) for erythroid differentiation and nearly doubled (0.444 $\rightarrow$ 0.799) in granulocyte-monocyte development, accompanied by a 0.228 absolute F1 score improvement.

Notably, KINDLE’s superiority proved most pronounced in AUPRC and F1 metrics. As summarized in Table 2, the validated edges in ground truth network constituting merely 0.9% (mESC), 0.65% (mHSC-E), 0.7% (mHSC-GM), and 1.15% (mHSC-L) of candidate edges, thus introducing severe class imbalance in the binary classification task. Under such conditions, AUPRC and F1 serve as more reliable performance indicators than AUROC (detailed justifications are provided in Appendix B.3). Therefore, the consistent AUPRC and F1 improvements demonstrated that privileged knowledge distillation provided a robust, prior-free route to GRN inference, outperforming not only expression-only algorithms but also methods that rely on explicit biological priors.

4.2 KINDLE Identified Key TFs and Their Stage-Specific Functions

Beyond quantitative metrics for assessing GRN accuracy, a crucial evaluation criterion lies in determining whether the inferred GRN can identify key TFs that orchestrate differentiation processes. SCENIC [20] introduced the AUCell algorithm, which calculates AUC scores for TF regulon (the set of all target genes of a TF in the inferred GRN) through rank-based enrichment analysis. This score reflects the functional activity of the TF regulon within each cell, enabling systematic identification

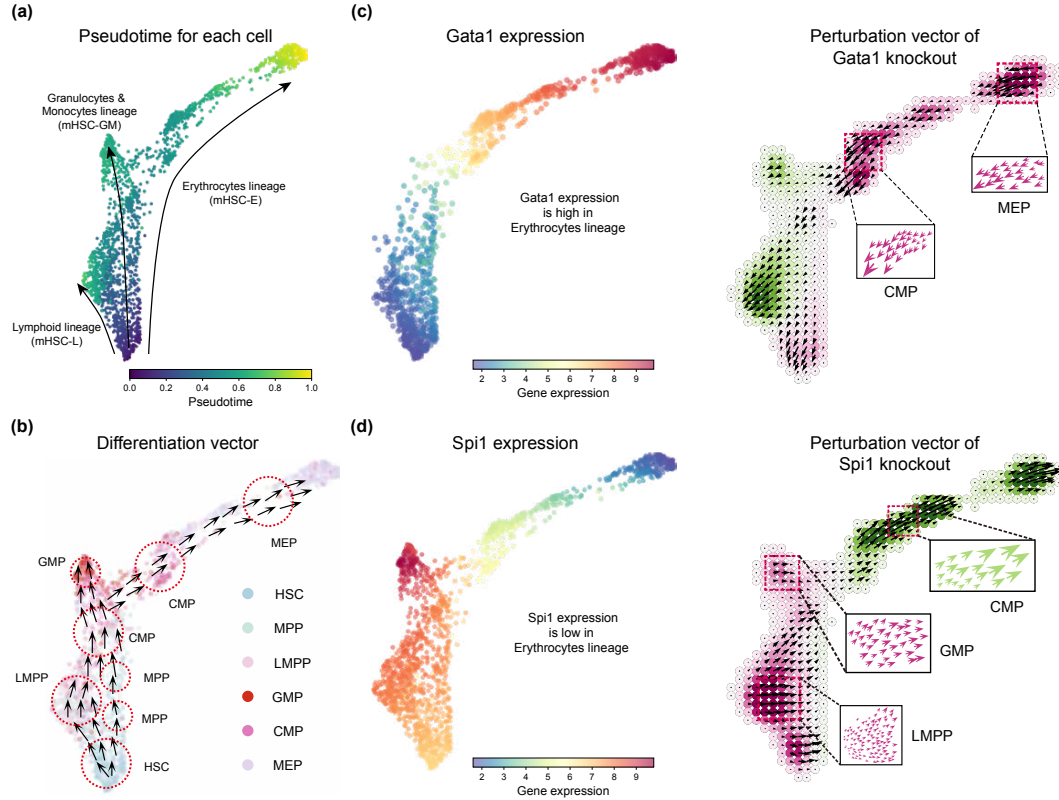


Figure 4: In silico perturbation analysis validates that KINDLE can accurately predict the cell fate transitions during the multi-lineage haematopoietic differentiation. (A) Slingshot-derived pseudotime trajectory embedding. (B) Cell-type annotations overlaid on the pseudotime landscape. (C) **Left:** Expression profiles of Gata1. **Right:** Silencing Gata1 induces reverse perturbation vectors along the erythroid branch and leads to differentiation arrest in CMP and MEP. (D) **Left:** Expression profiles of Spi1. **Right:** Silencing Spi1 represses the differentiation of LMPP and GMP and promotes erythroid progression in CMP.

of key regulators that drive cellular state transitions (see Appendix C for detailed methodology of AUCell). Following this protocol, we computed AUC scores from the GRN inferred by KINDLE-Gaussian to estimate per-cell TF activities in the mESC dataset. Given the five differentiation stages in this dataset (see Appendix A.1), we performed analysis of variance on AUC scores to assess whether they varied significantly throughout the differentiation stages and defined key TFs as those with Benjamini-Hochberg-adjusted P-values < 0.01 . We identified 25 key TFs, with their adjusted P-values listed in Table 3. Notably, 18 (72%) of 25 identified regulators have established roles in mESC differentiation according to prior literature [21–41].

Temporal patterning of the 25 TF regulon activities was investigated through hierarchical clustering of z-scored AUC scores (Figure 3). Two anti-correlated activation modules emerged from this unsupervised analysis: **Early-stage regulators** (Nanog, Sox2, Nr0b1, etc.) demonstrated peak activity at stage one with progressive attenuation through subsequent stages. This temporal trend aligns with known biological functions, such as Nanog and Sox2 being highly expressed in the early stage of mouse embryonic stem cells, maintaining the pluripotency of stem cells [42]. Their expression rapidly decreases as cells commit to differentiation, reflecting their pivotal function in regulating the transition from a pluripotent state to more specialized lineages [24]. **Late-stage regulators** (Gata4, Sox17, Kdm5b, etc.) exhibited minimal initial activity but showed significant activation from stage three onward. These results corroborate established mechanisms of lineage specification, where Sox17 overexpression upregulates a set of endoderm-specific gene markers and induces an ESC differentiation program towards primitive endoderm [43]. The emergence of these antiphasic expression patterns demonstrated that KINDLE not only recovered biologically relevant TFs but also assigned each regulator to its stage-specific functional context, thereby elucidating the sequential deployment of transcriptional programs during mESC differentiation.

4.3 KINDLE Predicted Lineage-Specific Fate Changes in In-Silico Perturbation

Following the precise identification of key TFs, a practical application involves interrogating their functional roles through systematic perturbation. We employed the Celloracle framework [16] to implement in silico perturbation analysis. This approach simulates TF knockout by setting target TF expression to zero and propagating the perturbation signal through the GRN’s topological structure to its target genes, ultimately generating a two-dimensional perturbation vector for each cell that predicts its fate trajectory under the specified perturbation (see Appendix D for algorithmic details). To validate the biological relevance of KINDLE-Gaussian’s predictions, we applied this methodology to the mHSC dataset, an ideal and complex benchmark system containing six distinct cell types (HSC, MPP, LMPP, GMP, CMP, MEP; see Appendix A.1 for cell type information) organized along three differentiation trajectories (Figure 4 a,b). The sequential differentiation order of different cell types is shown in Figure 5.

We focused on two well-characterized regulators governing hematopoietic lineage commitment, Gata1 and Spi1. Consistent with their established roles, Gata1 expression dominated in erythroid-lineage cells (CMP and MEP, Figure 4c), while Spi1 showed myeloid-lineage enrichment (LMPP and GMP, Figure 4d). Following the perturbation of Gata1, we generated perturbation vectors for each cell. Notably, the vectors for all cells within the erythroid lineage were opposite to the developmental direction of pseudotime trajectory shown in Figure 4a. This observation indicated that in the absence of Gata1, cells tend to revert to earlier progenitor states rather than progress towards more mature cell identities. To quantitatively assess the perturbation effect, we calculated a perturbation score for each cell and coloured the cells (purple $\rightarrow$ negative score, differentiation inhibited; green $\rightarrow$ positive score, differentiation promoted, see Appendix D.3 for additional details of perturbation score calculation). In erythroid lineage, all cells received negative scores, with the strongest inhibitory effect concentrated in CMP and MEP, the cell populations with the highest Gata1 expression levels. Subsequently, we applied the same perturbation procedure to Spi1. Upon silencing Spi1, all cells showed a developmental trajectory towards the erythroid lineage, with CMP differentiation being promoted as well as GMP and LMPP differentiation being inhibited (Figure 4d). These perturbation results are consistent with previous reports [44–47], where Gata1 promotes the differentiation of CMP into MEP (resulting in inhibition of CMP and MEP differentiation when Gata1 knockout), and Spi1 suppresses the CMP to MEP transition (resulting in CMP perturbation vectors pointing towards MEP upon Spi1 silencing).

Collectively, the in silico perturbation analyses demonstrated that within the haematopoietic system, KINDLE accurately modeled the downstream effects of key TFs knockout. Hence, beyond pinpointing key TFs, KINDLE provided a mechanistic scaffold for the rational design of cell-fate-engineering strategies.

4.4 Implementation Details

We trained KINDLE on an 80GB Nvidia A100 GPU for 30 epochs with a batch size of 32. To prevent overfitting, we implemented an early stopping strategy with a patience value set to 3. For optimization, we employed the Adam optimizer in conjunction with a warmup strategy, gradually increasing the learning rate from 0 to $1e-4$. Subsequently, a CosineAnnealingLR scheduler was utilized to further fine-tune the learning rate. During the training process, we conducted experiments with five distinct values for the hyperparameter W (1, 2, 4, 8, and 16), ultimately selecting $W = 16$ as it yielded the optimal results reported in this paper.

5 Discussion and Limitations

KINDLE advances GRN inference methodology by decoupling algorithm from prior knowledge dependency (the longstanding bottleneck in the field). Through integrating temporal causality modeling with knowledge distillation, our framework successfully transfers regulatory insights learned from privileged prior-augmented data to a prior-free student model, enabling KINDLE to achieve state-of-the-art performance across four benchmark datasets. The model’s ability to recover key TFs governing lineage specification validates its capacity to capture biologically interactions and the accurate prediction of knockout effects on hematopoietic cell fate transition underscores its potential for elucidating dynamic regulatory mechanisms in development and disease. The framework’s prior-independent nature positions it as a versatile tool for studying poorly characterized

systems, such as non-model organisms or emerging pathological states, where reliable prior networks are often unavailable.

Despite its strengths, KINDLE has several limitations. First, its reliance on temporal gene expression data restricts applicability to datasets with longitudinal sampling. Second, the distillation process may inherit biases from the teacher model’s prior-dependent training phase, potentially propagating errors from incomplete or noisy priors. Third, the current implementation focuses on transcriptional regulation, omitting post-transcriptional and epigenetic layers of gene regulation that could refine network predictions. Addressing these challenges will be critical for extending the framework’s utility across diverse biological contexts.

Acknowledgments and Disclosure of Funding

This research was supported by National Key Research and Development Program of China (2024YFF0507400) and National Natural Science Foundation of China (6220071694).

References

- [1] Pau Badia-i Mompel, Lorna Wessels, Sophia Müller-Dott, Rémi Trimbou, Ricardo O Ramirez Flores, Ricard Argelaguet, and Julio Saez-Rodriguez. Gene regulatory network inference in the era of single-cell multi-omics. *Nature Reviews Genetics*, 24(11):739–754, 2023.
- [2] Albert-Laszlo Barabasi and Zoltan N Oltvai. Network biology: understanding the cell’s functional organization. *Nature reviews genetics*, 5(2):101–113, 2004.
- [3] Guy Karlebach and Ron Shamir. Modelling and analysis of gene regulatory networks. *Nature reviews Molecular cell biology*, 9(10):770–780, 2008.
- [4] Jason D Buenrostro, Beijing Wu, Ulrike M Litzenburger, Dave Ruff, Michael L Gonzales, Michael P Snyder, Howard Y Chang, and William J Greenleaf. Single-cell chromatin accessibility reveals principles of regulatory variation. *Nature*, 523(7561):486–490, 2015.
- [5] Erez Lieberman-Aiden, Nynke L Van Berkum, Louise Williams, Maxim Imakaev, Tobias Ragoczy, Agnes Telling, Ido Amit, Bryan R Lajoie, Peter J Sabo, Michael O Dorschner, et al. Comprehensive mapping of long-range interactions reveals folding principles of the human genome. *science*, 326(5950):289–293, 2009.
- [6] Shuo Yang, Sujay Sanghavi, Holakou Rahmanian, Jan Bakus, and Vishwanathan SVN. Toward understanding privileged features distillation in learning-to-rank. *Advances in Neural Information Processing Systems*, 35:26658–26670, 2022.
- [7] Rui Peng, Juntian Qi, Yuxing Lu, Wei Wu, Qichen Sun, Chi Zhang, Yihan Chen, and Jinzhao Wang. Dissecting dynamic gene regulatory network using transformer-based temporal causality analysis. *bioRxiv*, pages 2025–02, 2025.
- [8] Thomas Moerman, Sara Aibar Santos, Carmen Bravo González-Blas, Jaak Simm, Yves Moreau, Jan Aerts, and Stein Aerts. Grnboost2 and arboreto: efficient and scalable inference of gene regulatory networks. *Bioinformatics*, 35(12):2159–2161, 2019.
- [9] Vân Anh Huynh-Thu, Alexandre Irrthum, Louis Wehenkel, and Pierre Geurts. Inferring regulatory networks from expression data using tree-based methods. *PloS one*, 5(9):e12776, 2010.
- [10] Alicia T Specht and Jun Li. Leap: constructing gene co-expression networks for single-cell rna-sequencing data using pseudotime ordering. *Bioinformatics*, 33(5):764–766, 2017.
- [11] Hirotaka Matsumoto, Hisanori Kiryu, Chikara Furusawa, Minoru SH Ko, Shigeru BH Ko, Norio Gouda, Tetsutaro Hayashi, and Itoshi Nikaido. Scode: an efficient regulatory network inference algorithm from single-cell rna-seq during differentiation. *Bioinformatics*, 33(15):2314–2321, 2017.

- [12] Thalia E Chan, Michael PH Stumpf, and Ann C Babbie. Gene regulatory network inference from single-cell data using multivariate information measures. *Cell systems*, 5(3):251–267, 2017.
- [13] Shahin Mohammadi, Jose Davila-Velderrain, and Manolis Kellis. Reconstruction of cell-type-specific interactomes at single-cell resolution. *Cell systems*, 9(6):559–568, 2019.
- [14] Qiuyue Yuan and Zhana Duren. Inferring gene regulatory networks from single-cell multiome data using atlas-scale external data. *Nature Biotechnology*, pages 1–11, 2024.
- [15] Peizhuo Wang, Xiao Wen, Han Li, Peng Lang, Shuya Li, Yipin Lei, Hantao Shu, Lin Gao, Dan Zhao, and Jianyang Zeng. Deciphering driver regulators of cell fate decisions from single-cell transcriptomics data with cefcon. *Nature Communications*, 14(1):8459, 2023.
- [16] Kenji Kamimoto, Blerta Stringa, Christy M Hoffmann, Kunal Jindal, Lilianna Solnica-Krezel, and Samantha A Morris. Dissecting cell identity via network inference and in silico gene perturbation. *Nature*, 614(7949):742–751, 2023.
- [17] Zhengrui Guo, Fangxu Zhou, Wei Wu, Qichen Sun, Lishuang Feng, Jinzhuo Wang, and Hao Chen. BLEND: Behavior-guided neural population dynamics modeling via privileged knowledge distillation. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [18] Aditya Pratapa, Amogh P Jalihal, Jeffrey N Law, Aditya Bharadwaj, and TM Murali. Benchmarking algorithms for gene regulatory network inference from single-cell transcriptomic data. *Nature methods*, 17(2):147–154, 2020.
- [19] Yijie Wang, Dong-Yeon Cho, Hangnoh Lee, Justin Fear, Brian Oliver, and Teresa M Przytycka. Reprogramming of regulatory network using expression uncovers sex-specific gene regulation in drosophila. *Nature communications*, 9(1):4061, 2018.
- [20] Sara Aibar, Carmen Bravo González-Blas, Thomas Moerman, Vân Anh Huynh-Thu, Hana Imrichova, Gert Hulselmans, Florian Rambow, Jean-Christophe Marine, Pierre Geurts, Jan Aerts, et al. Scenic: single-cell regulatory network inference and clustering. *Nature methods*, 14(11):1083–1086, 2017.
- [21] Claire Soudais, Malgorzata Bielinska, Markku Heikinheimo, Craig A MacArthur, Naoko Narita, Jeffrey E Saffitz, M Celeste Simon, Jeffrey M Leiden, and David B Wilson. Targeted mutagenesis of the transcription factor gata-4 gene in mouse embryonic stem cells disrupts visceral endoderm differentiation in vitro. *Development*, 121(11):3877–3888, 1995.
- [22] Benjamin L Kidder, Gangqing Hu, and Keji Zhao. Kdm5b focuses h3k4 methylation near promoters and enhancers during embryonic stem cell self-renewal and differentiation. *Genome biology*, 15:1–19, 2014.
- [23] Janel L Kopp, Briana D Ormsbee, Michelle Desler, and Angie Rizzino. Small increases in the level of sox2 trigger the differentiation of mouse embryonic stem cells. *Stem cells*, 26(4):903–911, 2008.
- [24] Yui-Han Loh, Qiang Wu, Joon-Lin Chew, Vinsensius B Vega, Weiwei Zhang, Xi Chen, Guillaume Bourque, Joshy George, Bernard Leong, Jun Liu, et al. The oct4 and nanog transcription network regulates pluripotency in mouse embryonic stem cells. *Nature genetics*, 38(4):431–440, 2006.
- [25] Guangjin Pan and James A Thomson. Nanog and transcriptional networks in embryonic stem cell pluripotency. *Cell research*, 17(1):42–49, 2007.
- [26] Akihiko Okuda, Akiko Fukushima, Masazumi Nishimoto, Akira Orimo, Toshiyuki Yamagishi, Yoko Nabeshima, Makoto Kuro-o, Yo-ichi Nabeshima, Kathy Boon, Marie Keaveney, et al. Utl1, a novel transcriptional coactivator expressed in pluripotent embryonic stem cells and extra-embryonic cells. *The EMBO journal*, 1998.

- [27] Fariba Dehghanian, Patrick Piero Bovio, Fabian Gather, Simone Probst, Amirhosein Naghsh-Nilchi, and Tanja Vogel. Zfp982 confers mouse embryonic stem cell characteristics by regulating expression of nanog, zfp42, and dppa3. *Biochimica et Biophysica Acta (BBA)-Molecular Cell Research*, 1871(4):119686, 2024.
- [28] Eiichi Okamura, Oliver H Tam, Eszter Posfai, Lingyu Li, Katie Cockburn, Cheryl QE Lee, Jodi Garner, and Janet Rossant. Esrrb function is required for proper primordial germ cell development in presomite stage mouse embryos. *Developmental biology*, 455(2):382–392, 2019.
- [29] Xiaofei Zhang, Juan Zhang, Tao Wang, Miguel A Esteban, and Duanqing Pei. Esrrb activates oct4 transcription and sustains self-renewal and pluripotency in embryonic stem cells. *Journal of Biological Chemistry*, 283(51):35825–35833, 2008.
- [30] Tianji Chen, Juan Du, and Guangxiu Lu. Cell growth arrest and apoptosis induced by oct4 or nanog knockdown in mouse embryonic stem cells: a possible role of trp53. *Molecular biology reports*, 39:1855–1861, 2012.
- [31] Setsuko Fujii, Satomi Nishikawa-Torikai, Yoko Futatsugi, Yayoi Toyooka, Mariko Yamane, Satoshi Ohtsuka, and Hitoshi Niwa. Nr0b1 is a negative regulator of zscan4c in mouse embryonic stem cells. *Scientific reports*, 5(1):9146, 2015.
- [32] Diego Pasini, Adrian P Bracken, Jacob B Hansen, Manuela Capillo, and Kristian Helin. The polycomb group protein suz12 is required for embryonic stem cell differentiation. *Molecular and cellular biology*, 2007.
- [33] Diego Pasini, Adrian P Bracken, Michael R Jensen, Eros Lazzerini Denchi, and Kristian Helin. Suz12 is essential for mouse development and for ezh2 histone methyltransferase activity. *The EMBO journal*, 23(20):4061–4071, 2004.
- [34] Kathy K Niakan, Hongkai Ji, René Maehr, Steven A Vokes, Kit T Rodolfa, Richard I Sherwood, Mariko Yamaki, John T Dimos, Alice E Chen, Douglas A Melton, et al. Sox17 promotes differentiation in mouse embryonic stem cells by directly regulating extraembryonic gene expression and indirectly antagonizing self-renewal. *Genes & development*, 24(3):312–326, 2010.
- [35] Ismail Kola, Sarah Brookes, Anthony R Green, Richard Garber, Martin Tymms, Takis S Papas, and Arun Seth. The ets1 transcription factor is widely expressed during murine embryo development and is associated with mesodermal cells involved in morphogenetic processes such as organ formation. *Proceedings of the National Academy of Sciences*, 90(16):7588–7592, 1993.
- [36] Sanjay K Singh, Mohamedi N Kagalwala, Jan Parker-Thornburg, Henry Adams, and Sadhan Majumder. Rest maintains self-renewal and pluripotency of embryonic stem cells. *Nature*, 453(7192):223–227, 2008.
- [37] Anna Połec, Alexander D Rowe, Pernille Blicher, Rajikala Suganthan, Magnar Bjørås, and Stig Ove Bøe. Pml regulates the epidermal differentiation complex and skin morphogenesis during mouse embryogenesis. *Genes*, 11(10):1130, 2020.
- [38] Yaser Atlasi, Rubina Noori, Claudia Gaspar, Patrick Franken, Andrea Sacchetti, Haleh Rafati, Tokameh Mahmoudi, Charles Decraene, George A Calin, Bradley J Merrill, et al. Wnt signaling regulates the lineage differentiation potential of mouse embryonic stem cells through tcf3 down-regulation. *PLoS genetics*, 9(5):e1003424, 2013.
- [39] Helen M Rowe, Adamandia Kapopoulou, Andrea Corsinotti, Liana Fasching, Todd S Macfarlan, Yara Tarabay, Stéphane Viville, Johan Jakobsson, Samuel L Pfaff, and Didier Trono. Trim28 repression of retrotransposon-based enhancers is necessary to preserve transcriptional dynamics in embryonic stem cells. *Genome research*, 23(3):452–461, 2013.
- [40] Eunjin Cho, Matthew R Mysliwiec, Clayton D Carlson, Aseem Ansari, Robert J Schwartz, and Youngsook Lee. Cardiac-specific developmental and epigenetic functions of jarid2 during embryonic development. *Journal of Biological Chemistry*, 293(30):11659–11673, 2018.

- [41] Songhwa Kang, Jisoo Yun, Seok Yun Jung, Yeon Ju Kim, Ji Hye Park, Seung Taek Ji, Woong Bi Jang, Jongseong Ha, Jae Ho Kim, Sang Hong Baek, et al. Adequate concentration of b cell leukemia/lymphoma 3 (bcl3) is required for pluripotency and self-renewal of mouse embryonic stem cells via downregulation of nanog transcription. *BMB reports*, 51(2):92, 2018.
- [42] Shinji Masui, Yuhki Nakatake, Yayoi Toyooka, Daisuke Shimosato, Rika Yagi, Kazue Takahashi, Hitoshi Okochi, Akihiko Okuda, Ryo Matoba, Alexei A Sharov, et al. Pluripotency governed by sox2 via regulation of oct3/4 expression in mouse embryonic stem cells. *Nature cell biology*, 9(6):625–635, 2007.
- [43] Xue-Bin Qu, Jie Pan, Cong Zhang, and Shu-Yang Huang. Sox17 facilitates the differentiation of mouse embryonic stem cells into primitive and definitive endoderm in vitro. *Development, growth & differentiation*, 50(7):585–593, 2008.
- [44] Kinuko Ohneda and Masayuki Yamamoto. Roles of hematopoietic transcription factors gata-1 and gata-2 in the development of red blood cell lineage. *Acta haematologica*, 108(4):237–245, 2002.
- [45] Yuko Fujiwara, Carol P Browne, Kerrianne Cunliffe, Sabra C Goff, and Stuart H Orkin. Arrested development of embryonic red cell precursors in mouse embryos lacking transcription factor gata-1. *Proceedings of the National Academy of Sciences*, 93(22):12355–12358, 1996.
- [46] Laura Gutiérrez, Noemí Caballero, Luis Fernández-Calleja, Elena Karkoulia, and John Strouboulis. Regulation of gata1 levels in erythropoiesis. *IUBMB life*, 72(1):89–105, 2020.
- [47] Pu Zhang, Xiaobo Zhang, Atsushi Iwama, Channing Yu, Kent A Smith, Beatrice U Mueller, Salaija Narravula, Bruce E Torbett, Stuart H Orkin, and Daniel G Tenen. Pu. 1 inhibits gata-1 function and erythroid differentiation by blocking gata-1 dna binding. *Blood, The Journal of the American Society of Hematology*, 96(8):2641–2648, 2000.
- [48] Kelly Street, Davide Risso, Russell B Fletcher, Diya Das, John Ngai, Nir Yosef, Elizabeth Purdom, and Sandrine Dudoit. Slingshot: cell lineage and pseudotime inference for single-cell transcriptomics. *BMC genomics*, 19:1–16, 2018.
- [49] Ronald R Coifman, Stephane Lafon, Ann B Lee, Mauro Maggioni, Boaz Nadler, Frederick Warner, and Steven W Zucker. Geometric diffusions as a tool for harmonic analysis and structure definition of data: Diffusion maps. *Proceedings of the national academy of sciences*, 102(21):7426–7431, 2005.
- [50] ENCODE Project Consortium et al. An integrated encyclopedia of dna elements in the human genome. *Nature*, 489(7414):57, 2012.
- [51] Zhaonan Zou, Tazro Ohta, and Shinya Oki. Chip-atlas 3.0: a data-mining suite to explore chromosome architecture together with large-scale regulome data. *Nucleic Acids Research*, 52(W1):W45–W53, 2024.
- [52] Huilei Xu, Caroline Baroukh, Ruth Dannenfelser, Edward Y Chen, Christopher M Tan, Yan Kou, Yujin E Kim, Ihor R Lemischka, and Avi Ma’ayan. Escape: database for integrating high-content published data collected from human and mouse embryonic stem cells. *Database*, 2013:bat045, 2013.
- [53] Luz Garcia-Alonso, Christian H Holland, Mahmoud M Ibrahim, Denes Turei, and Julio Saez-Rodriguez. Benchmark and integration of resources for the estimation of human transcription factor activities. *Genome research*, 29(8):1363–1375, 2019.
- [54] Zhi-Ping Liu, Canglin Wu, Hongyu Miao, and Hulin Wu. Regnetwork: an integrated database of transcriptional and post-transcriptional regulatory networks in human and mouse. *Database*, 2015:bav095, 2015.
- [55] Heonjong Han, Jae-Won Cho, Sangyoung Lee, Ayoung Yun, Hyojin Kim, Dasom Bae, Sunmo Yang, Chan Yeong Kim, Muyeon Lee, Eunbeen Kim, et al. Trustrust v2: an expanded reference database of human and mouse transcriptional regulatory interactions. *Nucleic acids research*, 46(D1):D380–D386, 2018.

- [56] Damian Szklarczyk, Annika L Gable, David Lyon, Alexander Junge, Stefan Wyder, Jaime Huerta-Cepas, Milan Simonovic, Nadezhda T Doncheva, John H Morris, Peer Bork, et al. String v11: protein–protein association networks with increased coverage, supporting functional discovery in genome-wide experimental datasets. *Nucleic acids research*, 47(D1):D607–D613, 2019.
- [57] Jesse Davis and Mark Goadrich. The relationship between precision-recall and roc curves. In *Proceedings of the 23rd international conference on Machine learning*, pages 233–240, 2006.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We provide experimental evidence which supports our claims and contributions in Section 4.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss the limitations in Section 5.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: This article is application-oriented and develops an algorithm to address the practical problems existing in the field of gene regulatory networks, rather than being a theoretical derivation article.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provided comprehensive details regarding all datasets and baseline methods in Appendix A. Additionally, Section 4.4 elaborates on the hardware infrastructure employed during training, along with specifics such as optimizer configuration, batch size selection, and other relevant hyperparameters.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: All datasets used in this study are publicly available, and detailed descriptions can be found in Appendix A.1. Additionally, the code implementation of KINDLE is thoroughly documented in Appendix F.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We provided comprehensive details about the hardware infrastructure employed during training, along with specifics such as optimizer configuration, batch size selection, and other relevant hyperparameters in Section 4.4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: Table 3 shows the P-values of the 25 transcription factors identified by KINDLE.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We elaborated on the compute resources during our training process in Section 4.4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research conducted conforms with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: In Section 5, we discussed the advantages and disadvantages of KINDLE, as well as its practical value and impact in areas lacking prior knowledge.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We do not released data or models that have a high risk for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We cited all the papers that are relevant to the code, models, and datasets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: We provided the full set of pseudocode for the implementation of KINDLE in Appendix F.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: This article does not involve any human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: This article does not involve any human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: Large language models are not important, original, or non-standard components of the core methods in this study.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Supplementary Contents of Datasets and Baselines

A.1 Datasets

Mouse embryonic stem cell (mESC). The mESC dataset contains Single-cell RNA sequencing (scRNA-seq) expression measurements for 421 primitive endoderm cells differentiated from mESCs, collected at five time points (0, 12, 24, 48, and 72 hours). Pseudotime computation was performed using Slingshot [48], with cells at 0 hours as the starting cluster and cells at 72 hours as the terminal differentiation state.

Mouse hematopoietic stem cell (mHSC). The mHSC dataset comprises 1,656 hematopoietic stem and progenitor cells traversing six critical differentiation states: hematopoietic stem cells (HSCs), multipotent progenitors (MPPs), lymphoid-primed multipotent progenitors (LMPPs), common myeloid progenitors (CMPs), megakaryocyte-erythrocyte progenitors (MEPs), and granulocyte-monocyte progenitors (GMPs). As visualized in Figure 5, these cell types follow distinct differentiation trajectories across three developmental lineages. Pseudotime trajectories were computed using the first three principal dimensions derived from DiffusionMap [49]. Gene regulatory networks were independently reconstructed for each lineage.

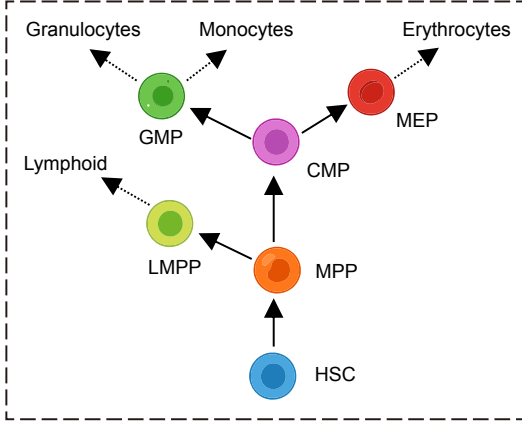


Figure 5: Differentiation schematic of six cell states in the mHSC dataset.

A.2 Ground Truth Networks

To benchmark inferred GRN, BEELINE [18] constructed ground truth networks from three stratified categories:

- **Cell-type-specific networks:**
 - Sourced from ENCODE [50], ChIP-Atlas [51], and ESCAPE [52] databases
 - Matched to the scRNA-seq dataset’s cell lineage
 - Included *loss-of-function* or *gain-of-function* perturbation data from ESCAPE
- **Non-specific networks:**
 - DoRothEA [53]: Integrated regulatory interactions filtered by confidence levels:
 - * Level A: high-confidence ChIP-seq data
 - * Level B: likely-confidence interactions
 - RegNetwork [54]: Genome-wide TF–gene and TF–TF interactions across human and mouse
 - TRRUST [55]: Manually curated TF–TG pairs from literature mining
- **Functional networks:**
 - Derived from STRING [56] protein interaction databases
 - Captured indirect regulatory effects (e.g., phosphorylation, co-expression)

Notably, in our study, cell-type-specific networks were employed as the ground truth for the benchmark evaluation of our model.

A.3 Baselines

GENIE3 [9]. GENIE3 decomposes GRN inference into p regression problems for p genes, using tree-based ensembles to quantify regulatory potential. For each target gene, it evaluates the predictive importance of all other genes’ expression patterns as putative regulators. These pairwise importance scores are aggregated to rank regulatory interactions and reconstruct directed networks.

GRNBoost2 [8]. GRNBoost2 is a gradient-boosting-based algorithm for GRN inference. Inspired by GENIE3, it trains tree-based regression models to predict each gene’s expression profile using TF expression data. The method employs regularized stochastic gradient boosting with an early-stopping heuristic: training terminates when out-of-bag data indicates non-improving loss function (average improvement < 0). Regulatory associations are aggregated and ranked by importance scores to construct the final GRN.

CEFCON [15]. CEFCON is a network control theory framework for cell fate analysis using scRNA-seq data. It constructs lineage-specific GRN via graph attention neural network under contrastive learning, aggregating gene interactions through adaptive neighborhood weighting. By integrating minimum feedback vertex sets and minimum dominating sets with GRN influence scores, it identifies driver regulators of cell fate transitions.

CellOracle [16]. CellOracle integrates scATAC-seq motif analysis and scRNA-seq data to model context-dependent GRNs. It first builds a base network through TF-binding motif scanning of regulatory DNA, and then refines edge weights using regularized linear models on expression data. Through in silico TF perturbation simulations, it predicts cell identity shifts by propagating signals across the GRN and analyzing pseudotime gradient vector fields.

NetREX [19]. NetREX reconstructs context-specific GRN by optimizing prior networks against expression data. Formulated as a non-convex l_0 -norm optimization problem, it iteratively modifies network topology using proximal alternative linearized maximization.

Random. The Random baseline generates GRN by randomly selecting k edges from all possible gene-gene interactions, where k equals the number of edges in the ground truth GRN.

Prior_Random. Prior_Random selects k edges exclusively from prior network interactions (k matches the number of ground truth GRN edges).

B Supplementary Contents of Benchmark Testing

B.1 The evaluation of GRN is regarded as a binary classification problem

The evaluation of inferred GRN can be formalized as a binary classification task, where edges in the inferred network are categorized relative to a ground truth GRN. As depicted in the Figure 6, let $\mathcal{G}_{\text{true}} = (V, E_{\text{true}})$ denote the ground truth network, and $\mathcal{G}_{\text{pred}} = (V, E_{\text{pred}})$ represent the inferred network. Each edge $e \in E_{\text{pred}}$ is classified into one of four categories: (1) True Positive (TP) : $e \in E_{\text{pred}} \cap E_{\text{true}}$. (2) False Positive (FP) : $e \in E_{\text{pred}} - E_{\text{true}}$. (3) True Negative (TN) : $e \notin E_{\text{pred}} \cup E_{\text{true}}$. (4) False Negative (FN) : $e \in E_{\text{true}} - E_{\text{pred}}$. Performance metrics are derived as follows: Precision = $\frac{TP}{TP+FP}$, Recall = $\frac{TP}{TP+FN}$, F1 = $2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$, TPR = $\frac{TP}{TP+FN}$, FPR = $\frac{FP}{FP+TN}$, AUROC = $\int_0^1 \text{TPR} d\text{FPR}$, AUPRC = $\int_0^1 \text{Precision} d\text{Recall}$. These metrics enable systematic comparison of GRN inference methods, quantifying both accuracy and robustness.

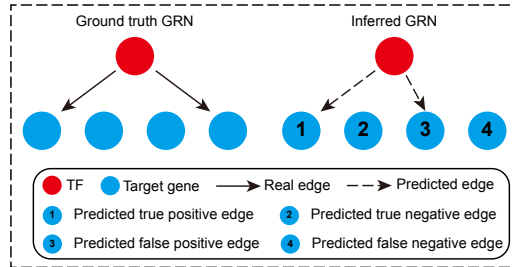


Figure 6: Schematic diagram of GRN evaluation, where the predicted edges can be classified into four types.

B.2 GRN Benchmarking Algorithm

Algorithm 1 GRN Benchmark Evaluation

```

1: Input:
2: - Ground truth GRN:  $\mathcal{G}_{\text{true}} = (V, E_{\text{true}})$ 
3: - Predicted GRN for  $N$  algorithms:  $\{\mathcal{G}_{\text{pred}}^{(i)} = (V, E_{\text{pred}}^{(i)})\}_{i=1}^N$ 
4: Output: Metrics  $\in \mathbb{R}^{N \times 3}$  ▷ DataFrame containing AUROC, AUPRC, F1
5: // Preprocess ground truth
6:  $E_{\text{true}} \leftarrow E_{\text{true}} - \{(v, v) \mid v \in V\}$  ▷ Remove self-loops
7:  $E_{\text{true}} \leftarrow \text{Deduplicate}(E_{\text{true}})$  ▷ Remove duplicate edges
8: results  $\leftarrow \emptyset$  ▷ Initialize metric collection
9: for each predicted GRN  $\mathcal{G}_{\text{pred}}^{(i)}$  do
10:    $E_{\text{pred}}^{(i)} \leftarrow \text{Sort}(\text{Deduplicate}(E_{\text{pred}}^{(i)}))$  ▷ Sort predicted edges
11:   // Generate candidate edges
12:    $P_{\text{all}} \leftarrow V \times V - \{(v, v)\}$  ▷ All potential non-self edges
13:    $\mathbf{y}_{\text{true}} \leftarrow [\mathbb{I}(e \in E_{\text{true}}) \mid \forall e \in P_{\text{all}}]$  ▷ Obtain true labels, where  $\mathbb{I}$  is the indicator function
14:   // Obtain predicted edges
15:    $\mathbf{y}_{\text{pred}}^{(i)} \leftarrow \text{TopK}(E_{\text{pred}}^{(i)})$  ▷ Select the top k edges
16:   // Calculate metrics
17:    $\text{TPR}^{(i)}, \text{FPR}^{(i)} \leftarrow \frac{\text{TP}}{\text{TP} + \text{FN}}, \frac{\text{FP}}{\text{FP} + \text{TN}}$  ▷ ROC components
18:    $\text{Precision}^{(i)}, \text{Recall}^{(i)} \leftarrow \frac{\text{TP}}{\text{TP} + \text{FP}}, \frac{\text{TP}}{\text{TP} + \text{FN}}$  ▷ PR components
19:    $\text{F1}^{(i)} \leftarrow 2 \cdot \frac{\text{Precision}^{(i)} \cdot \text{Recall}^{(i)}}{\text{Precision}^{(i)} + \text{Recall}^{(i)}}$  ▷ Calculate F1 score
20:    $\text{AUROC}^{(i)}, \text{AUPRC}^{(i)} \leftarrow \int_0^1 \text{TPR} \, d\text{FPR}, \int_0^1 \text{Precision} \, d\text{Recall}$ 
21:   results  $\leftarrow \text{results} \cup \{(i, \text{AUROC}^{(i)}, \text{AUPRC}^{(i)}, \text{F1}^{(i)})\}$  ▷ Record metrics
22: end for
23: // Aggregate results
24: Metrics  $\leftarrow \text{ConstructDataFrame}(\text{results})$  ▷ Shape:  $N \times 3$ 
25: return Metrics

```

B.3 AUPRC is more robust compared to AUROC

Table 2: The distribution of positive and negative labels in four datasets

	mESC	mHSC-E	mHSC-GM	mHSC-L
Genes	1652	1933	1520	640
Potential edges	2727452	3734556	230880	408960
True edges	24557	24726	16198	4705
Proportion of true edges	0.90%	0.65%	0.70%	1.15%

The evaluation framework described in Appendix B.2 operates on a search space of TF–TG pairs defined as $\mathcal{E}_{\text{potential}} = M \times (M - 1)$, where M denotes the number of genes in a dataset. Within this space, edges present in the ground truth GRN are defined as $\mathcal{E}_{\text{true}}$. We quantified the distribution of $\mathcal{E}_{\text{potential}}$ and $\mathcal{E}_{\text{true}}$ in Table 2 and found that there is an extreme class imbalance inherent to GRN inference across four datasets. For instance:

$$\begin{aligned}
\text{mESC:} \quad & |\mathcal{E}_{\text{potential}}| = 2,727,452, \quad |\mathcal{E}_{\text{true}}| = 24,557 \quad (\sim 0.9\% \text{ positivity rate}) \\
\text{mHSC-L:} \quad & |\mathcal{E}_{\text{potential}}| = 408,960, \quad |\mathcal{E}_{\text{true}}| = 4,705 \quad (\sim 1.1\% \text{ positivity rate})
\end{aligned}$$

In such scenarios, AUROC disproportionately emphasizes the majority class (negative edges) due to its reliance on the false positive rate ($\text{FPR} = \frac{\text{FP}}{\text{FP} + \text{TN}}$). When $|\mathcal{E}_{\text{true}}| \ll |\mathcal{E}_{\text{potential}}|$, the sum $\text{FP} + \text{TN} \approx |\mathcal{E}_{\text{potential}}|$, which makes AUROC overly optimistic in evaluating model performance. Conversely, AUPRC is better equipped to handle such imbalanced scenarios [57]. As a result, AUPRC can more accurately reflect the performance of the model.

C Supplementary Contents of AUCell Algorithm

AUCell is designed to quantify the activity of predefined gene regulatory regulons in scRNA-seq data. By calculating the Area Under the Recovery Curve (AUC) for regulons across individual cells, it identifies cells exhibiting coordinated activation of specific transcriptional programs, independent of absolute expression scales.

For a regulon R comprising m genes and a cell c with n detected genes, AUCell operates through three sequential phases. First, genes in cell c are ranked by their expression values in descending order, generating an ordered list $g^c = (g_1^c, g_2^c, \dots, g_n^c)$, where g_1^c denotes the highest-expressed gene. Ties in expression values are resolved stochastically to avoid rank bias. Subsequently, a binary recovery vector is constructed using an indicator function for regulon membership:

$$\mathbb{I}_R(g_i^c) = \begin{cases} 1, & \text{if } g_i^c \in R \\ 0, & \text{otherwise} \end{cases} \quad (13)$$

The cumulative recovery score $S(R, k)$ is computed by summing $\mathbb{I}_R(g_i^c)$ over the top k genes:

$$S(R, k) = \sum_{i=1}^k \mathbb{I}_R(g_i^c) \quad (14)$$

By taking the gene ranking k as the x-axis and the cumulative recovery score $S(R, k)$ as the y-axis, a recovery curve can be plotted, as illustrated in Figure 7. Finally, the AUC score is obtained by calculating the area under this curve. In essence, the AUC score evaluates whether a crucial subset of the input gene set is preferentially enriched among the top-ranked genes in each cell. It also quantifies the proportion of expressed signature genes and their relative expression levels compared to all other genes within the cell, thereby providing a measure of the regulon’s activity in that specific cell.

In KINDLE, after calculating the AUC score, we additionally conducted an analysis of variance on this score. Based on statistical significance analysis, we ultimately identified 25 key TFs, with their corresponding P-values summarized in Table 3.

Table 3: P-values and adjusted P-values for selected TFs, with those highlighted in red corresponding to TFs previously reported in the literature.

Gene	P-value	adjusted P-value	Gene	P-value	adjusted P-value
Gata4	3.8203×10^{-16}	4.4315×10^{-16}	Sox17	2.4402×10^{-60}	3.9314×10^{-60}
Kdm5b	8.1186×10^{-58}	1.2392×10^{-57}	Mybl2	1.4571×10^{-43}	1.8372×10^{-43}
Sox2	1.2096×10^{-47}	1.6705×10^{-47}	Suz12	1.9448×10^{-19}	2.3500×10^{-19}
Nanog	2.5669×10^{-80}	6.7672×10^{-80}	Ets1	4.7752×10^{-73}	9.2321×10^{-73}
Utf1	2.4185×10^{-71}	4.3836×10^{-71}	Rest	6.1475×10^{-83}	1.9808×10^{-82}
Zfp42	7.8886×10^{-113}	4.5754×10^{-112}	Egr1	1.8206×10^{-82}	5.2797×10^{-82}
Esrrb	2.3245×10^{-108}	1.1235×10^{-107}	Pml	6.8755×10^{-80}	1.6616×10^{-79}
Trp53	3.8236×10^{-171}	1.1088×10^{-169}	Tcf7l2	1.2169×10^{-56}	1.7646×10^{-56}
Sap30	3.0508×10^{-119}	4.4236×10^{-118}	Trim28	1.8347×10^{-105}	7.6010×10^{-105}
Nr0b1	1.2291×10^{-78}	2.7419×10^{-78}	Elf1	1.9730×10^{-114}	1.4304×10^{-113}
NfyA	2.5722×10^{-47}	3.3907×10^{-47}	Jarid2	5.9858×10^{-68}	1.0211×10^{-67}
Tcf12	9.1998×10^{-119}	8.8931×10^{-118}	Bcl3	3.2749×10^{-73}	6.7838×10^{-73}
Akt1	1.3803×10^{-90}	5.0038×10^{-90}			

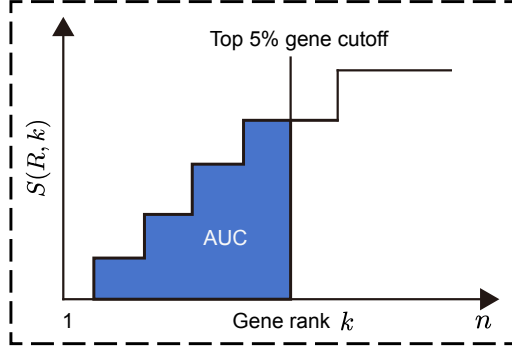


Figure 7: The recovery curve in the AUCell algorithm, with the gene ranking as x-axis and the recovery score as the y-axis.

D Supplementary Contents of In Silico Perturbation

In silico perturbation serves as a critical benchmark for evaluating the accuracy of GRN. By simulating TF perturbation (e.g., knockouts) and propagating their effects through the inferred GRN, this approach quantifies the network’s ability to predict downstream gene expression changes and cell fate transitions. The following sections detail the computational framework of CellOracle’s in silico perturbation pipeline, which integrates GRN-based signal propagation, cell-state transition modeling, and perturbation score calculation.

D.1 Signal Propagation for TF Perturbation Simulation

Given an inferred GRN represented by its regulatory coefficient matrix $\mathbf{A} \in \mathbb{R}^{M \times M}$, where M denotes the number of genes. When perturbing a TF i , we set its expression to zero. By subtracting the perturbed expression values from the original ones, we obtain a vector $\Delta \mathbf{x}^{(0)} \in \mathbb{R}^M$, which is defined as:

$$\Delta x_j^{(0)} = \begin{cases} -x_i, & \text{if } j = i \text{ (TF knockout)} \\ 0, & \text{otherwise} \end{cases} \quad (15)$$

The impact of this perturbation on gene expression is propagated through the matrix $\mathbf{A}$. For the first order perturbation, it is calculated as:

$$\Delta \mathbf{x}^{(1)} = \mathbf{A}^\top \Delta \mathbf{x}^{(0)} \quad (16)$$

where $\mathbf{A}^\top$, based on the regulatory weights between gene pairs, propagates the perturbation effect to their direct targets. Higher order indirect effects are computed iteratively via K rounds of signal propagation. Specifically:

$$\Delta \mathbf{x}^{(k)} = \mathbf{A}^\top \Delta \mathbf{x}^{(k-1)}, \quad k = 2, 3, \dots, K \quad (K = 3 \text{ for default}) \quad (17)$$

After the k -th propagation, the resulting perturbation vector $\Delta \mathbf{x}^{(k)}$ is considered as the simulated perturbation vector $\Delta \mathbf{X}_{\text{sim}}$. It should be noted that during each propagation step, if any element in $\Delta \mathbf{x}^{(k)}$ is less than 0, this element needs to be reassigned as 0, since gene expression levels are always non-negative. Mathematically, this can be expressed as:

$$\Delta \mathbf{x}^{(k)} \leftarrow \max(\Delta \mathbf{x}^{(k)}, 0) \quad (18)$$

D.2 Cell-State Transition Estimation

The simulated gene expression shifts $\Delta \mathbf{X}_{\text{sim}}$ are translated into cell-state transition probabilities through a kernelized similarity analysis in the two-dimensional embedding space. For each cell i , the transition probability $p_{i,j}$ to its K -nearest neighbors ($j \in \mathcal{N}_i$) is computed by comparing the simulated perturbation vector $\Delta \mathbf{X}_{\text{sim},i}$ with the observed expression difference $\mathbf{X}_j - \mathbf{X}_i$. This is formalized using a softmax function over Pearson correlation similarities:

$$p_{i,j} = \frac{\exp(\rho(\Delta \mathbf{X}_{\text{sim},i}, \mathbf{X}_j - \mathbf{X}_i)/\tau)}{\sum_{k \in \mathcal{N}_i} \exp(\rho(\Delta \mathbf{X}_{\text{sim},i}, \mathbf{X}_k - \mathbf{X}_i)/\tau)} \quad (19)$$

where ρ denotes the Pearson correlation function, $\mathbf{X}_j$ means the expression of gene j , $\mathbf{X}_i$ means the expression of gene i and $\tau = 0.05$ modulates the selectivity of the probability distribution. The transition probabilities are then projected onto the two-dimensional embedding space to construct a perturbation vector field. For each cell-neighbor pair, the coordinate difference vector $\mathbf{v}_{i,j} = \mathbf{V}_j - \mathbf{V}_i$ ($\mathbf{V}_i$ means the coordinate of gene i in the two-dimensional space) is weighted by $p_{i,j}$, yielding the simulated perturbation vector for cell i :

$$\mathbf{v}_{\text{sim},i} = \sum_{j \in \mathcal{N}_i} p_{i,j} \cdot \mathbf{v}_{i,j} \quad (20)$$

This vector $\mathbf{v}_{\text{sim},i}$ represents the predicted direction and magnitude of cell-state transition induced by the TF perturbation. To account for cellular heterogeneity, this process is repeated across all cells, generating a global perturbation vector field $\mathbf{V}_{\text{sim}} \in \mathbb{R}^{N \times 2}$, where N is the number of cells. This vector field captures context-dependent regulatory effects, enabling systematic visualization of simulated differentiation trajectories.

D.3 Perturbation Score Calculation

The perturbation score (PS) quantifies the alignment between simulated perturbation-driven cell-state transitions and intrinsic differentiation trajectories. The intrinsic differentiation vector field $\mathbf{V}_{\text{diff}} \in \mathbb{R}^{N \times 2}$ is derived as the spatial gradient of pseudotime t , where pseudotime (inferred via diffusion pseudotime or RNA velocity) represents the progression of cells along developmental trajectories. Specifically, $\mathbf{V}_{\text{diff},i} = \nabla t_i$ captures the direction and magnitude of natural differentiation for cell i in the low-dimensional embedding space. To evaluate the impact of TF perturbation, we compute the cosine similarity between the simulated perturbation vector $\mathbf{v}_{\text{sim},i}$ and the differentiation vector $\mathbf{v}_{\text{diff},i}$:

$$\text{PS}_i = \frac{\mathbf{v}_{\text{sim},i} \cdot \mathbf{v}_{\text{diff},i}}{\|\mathbf{v}_{\text{sim},i}\| \|\mathbf{v}_{\text{diff},i}\|} \quad (21)$$

where a positive PS (green in Figure 4c,d) indicates that the perturbation promotes differentiation along the native trajectory, while a negative PS (purple in Figure 4c,d) suggests suppression of differentiation. This directional alignment metric enables systematic identification of TFs that act as drivers or brakes in cell fate determination.

E Supplementary Contents of Whole TF's AUC Score

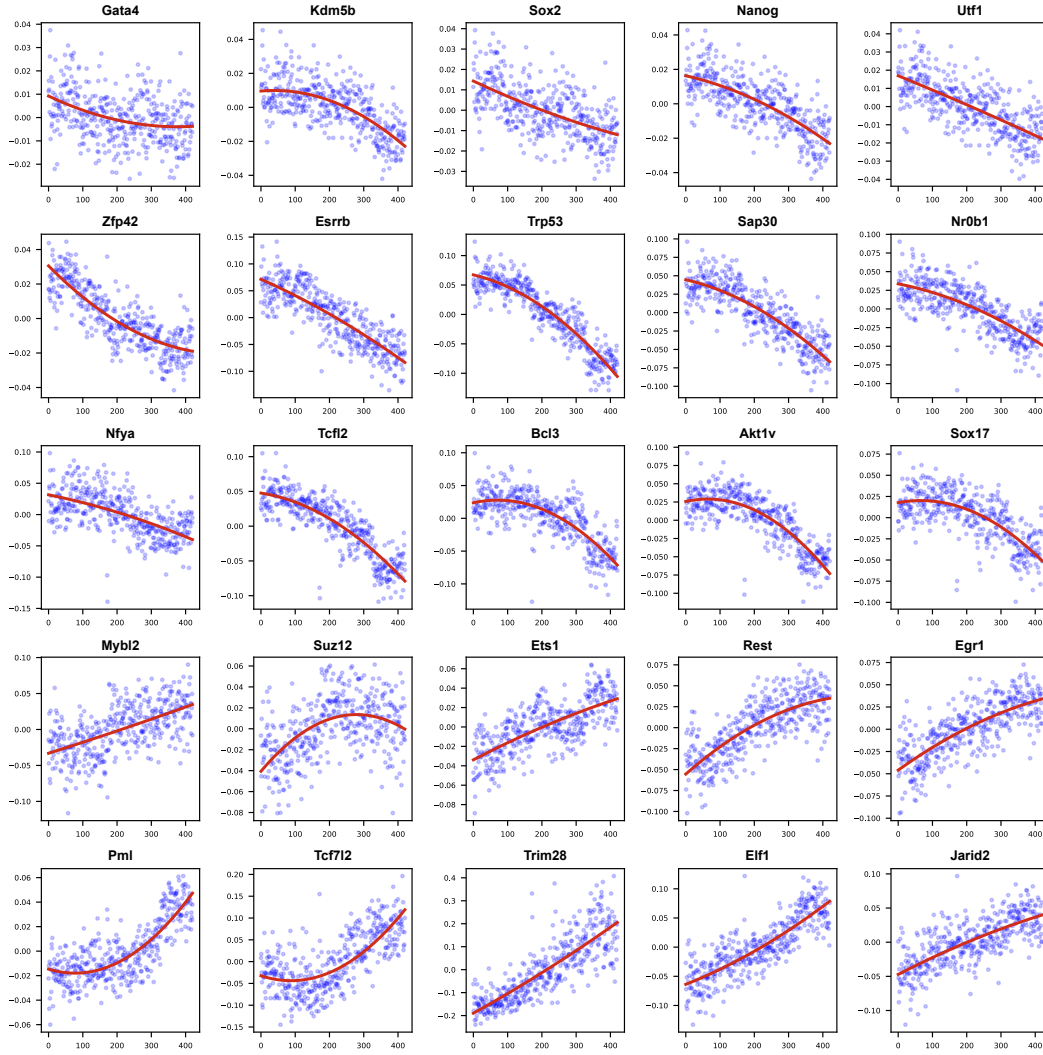


Figure 8: Temporal dynamics of AUC scores for all 25 TFs identified by KINDLE. Blue scatter points represent AUC scores at individual time points, while the red curve denotes a quadratic fitting of these scores.

F Supplementary Contents of KINDLE Algorithm

Algorithm 2 KINDLE Framework for GRN Inference

```

1: Input:
2: - Temporal expression matrix:  $\mathbf{G} \in \mathbb{R}^{N \times M}$  with  $N$  time points and  $M$  genes
3: - Spatial prior mask for teacher model:  $\mathcal{M}_{spatial} \in \{0, 1\}^{M \times M}$ 
4: Output: Inferred GRN adjacency matrix  $\mathcal{G}_{pred} \in \mathbb{R}^{M \times M}$ 
5: // Teacher Training Stage, Initialize Teacher  $f_{\theta_T}$  with parameters  $\Theta_T$ 
6: for epoch  $\in [1, E_{max}]$  do
7:   Sample batch  $\mathcal{B} \sim \mathbf{G}$  where  $|\mathcal{B}| = B$ 
8:   for each sequence  $X_{1:T+W}^{(i)}$  do ▷  $T$ : Historical window,  $W$ : Prediction window
9:     Partition  $X^{(i)} \rightarrow (X_{1:T}^{(i)}, Y_{T+1:T+W}^{(i)})$  ▷ Partition sequence into historical and future parts
10:     $Q_t, K_t, V_t = QKV(X_{1:T}^{(i)})$  ▷ Get  $Q, K, V$ 
11:     $A_t = \text{softmax}\left(\frac{Q_t K_t^\top}{\sqrt{d_k}} \odot \mathcal{M}_{temporal}\right)$  ▷ Compute temporal attention
12:     $H_t = A_t V_t \in \mathbb{R}^{B \times T \times M}$  ▷ Obtain input for spatial layer
13:     $Q_s, K_s, V_s = QKV(\psi(H_t))$  ▷  $\psi$ : Transposition operation
14:     $A_s = \text{softmax}\left(\frac{Q_s K_s^\top}{\sqrt{d_k}} \odot \mathcal{M}_{spatial}\right)$ 
15:     $\hat{Y}^{(i)} = \psi(A_s V_s) W_h \in \mathbb{R}^{B \times W \times M}$  ▷  $W_h$ : Weights for output layer
16:    Update  $\Theta_T \propto \nabla_{\Theta_T} \left[ \frac{1}{BWM} \sum \|\hat{Y}^{(i)} - Y^{(i)}\|_2^2 \right]$  ▷ Update teacher model parameters
17:   end for
18: end for
19: return Teacher model  $f_T(\cdot; \Theta_T)$ 
20: // Student Distillation Stage, Initialize Student  $f_{\theta_S}$  with parameters  $\Theta_S$  and frozen Teacher  $f_{\theta_T}$ 
21: for epoch  $\in [1, E_{distill}]$  do
22:   for each  $X_{1:T+W}^{(i)} \in \mathcal{B}$  do
23:     Partition  $X^{(i)} \rightarrow (X_{1:T}^{(i)}, Y_{T+1:T+W}^{(i)})$ 
24:      $\hat{Y}_T^{(i)} = f_T(X_{1:T}^{(i)})$  ▷ Get teacher predictions
25:      $\hat{Y}_S^{(i)} = f_S(X_{1:T}^{(i)})$  ▷ Get student predictions
26:      $\mathcal{L}_{pred} = \frac{1}{BWM} \sum \|\hat{Y}_S^{(i)} - Y^{(i)}\|_2^2$  ▷ Compute student prediction loss
27:     if Distillation Type = "Hard" then ▷ Choose distillation loss
28:        $\mathcal{L}_{distill} = \frac{1}{BWM} \sum \|\hat{Y}_T^{(i)} - \hat{Y}_S^{(i)}\|_2^2$ 
29:     else if Distillation Type = "Soft" then
30:        $\mathcal{L}_{distill} = \frac{1}{BWM} \sum \text{KL}\left(\sigma\left(\frac{\hat{Y}_T^{(i)}}{\tau}\right) \parallel \sigma\left(\frac{\hat{Y}_S^{(i)}}{\tau}\right)\right)$  ▷  $\tau > 0$ : Softmax temperature
31:     else if Distillation Type = "Bilinear" then
32:        $\mathcal{L}_{distill} = \frac{1}{BWM} \sum (\hat{Y}_T^{(i)})^\top (\hat{Y}_S^{(i)})$ 
33:     else if Distillation Type = "Gaussian" then
34:        $\mathcal{L}_{distill} = \frac{1}{BWM} \sum \exp\left(-\frac{\|\hat{Y}_T^{(i)} - \hat{Y}_S^{(i)}\|_2^2}{2\lambda^2}\right)$ 
35:     end if
36:      $\mathcal{L}_{total} = \alpha \mathcal{L}_{pred} + (1 - \alpha) \mathcal{L}_{distill}$  ▷  $\alpha \in [0, 1]$ : Knowledge distillation coefficient
37:     Update  $\Theta_S \propto \nabla_{\Theta_S} \mathcal{L}_{total}$  ▷ Update student model parameters
38:   end for
39: end for
40: return Student model  $f_S(\cdot; \Theta_S)$ 
41: // GRN Inference Stage
42: Partition  $\mathbf{G}$  into  $H$  subsequences  $\{\mathcal{S}^{(h)}\}_{h=1}^H$  with stride  $T$ 
43: Initialize adjacency confidence matrix  $\bar{A} = \mathbf{0}^{M \times M}$ 
44: for each  $\mathcal{S}^{(h)} \in \{\mathcal{S}^{(h)}\}_{h=1}^H$  do
45:    $\bar{A} \leftarrow \bar{A} + \frac{1}{H} \phi(f_S(\mathcal{S}^{(h)}))$  ▷  $\phi$ : Extract attention matrix in student's spatial layer
46: end for
47: return  $\mathcal{G}_{pred} = \text{top}_k(\bar{A})$  ▷ Select top k edges

```
