

A079 · GlintID: Interpretable Modeling of Combinatorial Regulatory Logic

Presenter: Arush Ramteke — New York University

Authors: Arush Ramteke, Simon Liu, Oded Regev

Session: Day 1 · Thu Oct 1, 2026 · 4:15–5:15 PM

Understanding the combinatorial logic by which DNA and RNA sequences regulate biological function is a central challenge in molecular biology. Deep learning models for sequence-to-function prediction achieve strong predictive performance. Their learned regulatory logic is typically interpreted via post-hoc methods that discover motifs from attribution scores and examine interactions between motifs by perturbing them or evaluating counterfactual motif combinations embedded in background sequences. However, prior work has shown that such analyses often provide inconsistent or unfaithful representations of model behavior. Importantly, they offer an incomplete view of the model's learned mechanism. Here, we introduce GlintID, an interpretable-by-design architecture that jointly discovers regulatory motifs and composes predictions from their individual and distance-dependent combinatorial effects, providing a complete quantitative view of the model's learned regulatory logic.

We applied GlintID to three sequence-to-function tasks: proximal versus distal polyadenylation site usage prediction, *Drosophila* enhancer activity prediction, and exon inclusion prediction (splicing). In all cases, GlintID approaches the performance of strong black-box architectures while outperforming an additive-effects-only ablation, demonstrating that explicit interaction modeling captures substantial predictive signals. Interpreting trained models reveals potentially novel motifs and interactions while also recapitulating regulatory syntax previously identified through post-hoc analyses. Altogether, our results establish interpretable-by-design modeling as a broadly applicable approach for systematically extracting experimentally testable regulatory hypotheses from sequence-to-function models.

GlntID: Interpretable Modeling of Combinatorial Regulatory Logic

Arush Ramteke^{*1}, Simon Liu¹, Oded Regev¹

¹ Courant Institute, New York University

Introduction. Regulatory function emerges not only from individual DNA and RNA sequence elements, or motifs, but also from the interactions between them. For example, they may reflect spacing- and orientation- dependent assembly of molecular complexes, or pioneering effects in which one element modulates the accessibility of another. Characterizing these interactions is essential for understanding the mechanisms driving gene regulation.

Deep learning models for sequence-to-function prediction have achieved strong performance across diverse regulatory processes. Post-hoc interpretation methods are then used to decipher their learned logic, including both motifs and interactions between them. For example, interactions can be inferred by perturbing motifs or evaluating counterfactual motif combinations embedded in a background sequence. However, prior work has shown that such analyses often provide inconsistent, unfaithful, or incomplete representations of model behavior [1,2]. As a result, these methods may fail to discover critical logic and even generate misleading results.

In contrast, *interpretable-by-design* models are constructed from separable components whose logic can be directly examined and visualized. As a result, such models provide a complete view of their learned mechanisms. Here, we introduce GlntID, a novel interpretable-by-design architecture that jointly discovers regulatory motifs and quantifies their individual and distance-dependent combinatorial effects.

Methods. GlntID composes each prediction as an additive combination of individual motif effects and pairwise interactions between learned regulatory motifs. For each sequence, convolutional filters identify F regulatory motifs and produce activation maps $A^{(i)}$, from which the locations of the top- k activations, T_i , are selected. The resulting contribution score is:

$$z = \sum_{i=1}^F \sum_{s \in T_i} f_i(A_s^{(i)}, s) + \sum_{1 \leq i < j \leq F} \sum_{s \in T_i} \sum_{t \in T_j} f_{ij}(A_s^{(i)} \cdot A_t^{(j)}, t - s).$$

Here, $A_s^{(i)}$ denotes the activation strength of motif i at position s . The functions f_i quantify individual motif effects, while f_{ij} quantify pairwise effects as a function of motif co-occurrence and relative spacing. Each function is parameterized by a neural network, enabling expressiveness while preserving an explicit decomposition of model predictions. Finally, a learnable tuner maps the summed contribution score to the prediction, $\hat{y} = g(z)$, capturing global nonlinearities such as saturation.

Results. We evaluated GlntID on three sequence-to-function tasks spanning polyadenylation site prediction, enhancer activity prediction, and exon

inclusion prediction (splicing) [3–5]. Across tasks, GlntID approaches the predictive performance of the best black-box architectures while outperforming an additive-effects-only ablation ($R^2 = 0.911/0.534/0.846$ for GlntID, $0.951/0.592/0.871$ for the best black-box model and $0.836/0.488/0.812$ for GlntID-Additive respectively), demonstrating that explicit interaction modeling captures significant predictive signals.

GlntID recapitulates regulatory logic previously extracted via post-hoc analyses of black-box models while revealing additional motifs and interactions. On the polyadenylation task, it learns a sparse set of 15 informative motifs spanning the core polyadenylation machinery and reconstructs a position-dependent regulatory network in which upstream and downstream regulatory elements interact with the canonical cleavage signal to coordinate polyadenylation (Fig. 1) [3]. Across the enhancer and splicing tasks, GlntID similarly recovers established regulatory mechanisms while identifying additional interactions and regulatory hypotheses supported by literature (not shown here).

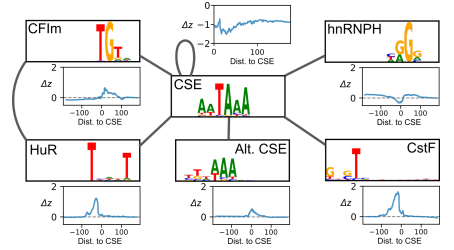


Figure 1. Subset of regulatory interaction network inferred by GlntID on polyadenylation site prediction task. Various motifs interact with the central sequence element (CSE). Contributions of each edge (interaction) shown as a function of the distance between the motif and the CSE.

Conclusion. GlntID learns a compact regulatory model that directly generates experimentally testable hypotheses. Its interpretable pairwise interactions capture much of the predictive signal learned by state-of-the-art black-box models. Incorporating higher-order interactions and profile prediction may potentially extend both model performance and discoverable regulatory logic.

References.

1. Zhou, S. *et al.*, *Proc. of 43rd ICML (PMLR)*, (2026).
2. Nagai, M. *et al.* *Nat. Genet.* **58**, (2026).
3. Bogard, N. *et al.* *Cell* **178**, (2019).
4. de Almeida, B. P. *et al.* *Nat. Genet.* **54**, (2022).
5. Liao, S. E. *et al.* *Proc. Natl. Acad. Sci.* **120**, (2023).