



MUGO: Differentiable Combinatorial Optimization for Causal Variant Discovery in the Non-coding Genome

Simon Dongbo Sun
University of California, Irvine
Irvine, CA, USA
dongbos@uci.edu

Yaqi Hu
University of California, Irvine
Irvine, CA, USA
yaqih10@uci.edu

Junhao Liu
University of California, Irvine
Irvine, CA, USA
junhaoliu17@gmail.com

Martin Jinye Zhang
Computational Biology
Carnegie Mellon University
Pittsburgh, Pennsylvania, USA
martinzh@andrew.cmu.edu

Pengcheng Xu
University of California, Irvine
Irvine, CA, USA
pengcheng.xu@uci.edu

Jing Zhang*
Donald Bren School of Information
and Computer Sciences, University of
California, Irvine
Irvine, CA, USA
Mathematical, Computational and
Systems Biology
University of California, Irvine
Irvine, CA, USA
zhang.jing@uci.edu

Abstract

Deciphering how non-coding variants perturb gene regulation is central to translating GWAS loci into mechanism, yet existing prioritization methods rarely deliver cell-type-resolved molecular effects, causal variant-to-gene attribution, or principled reasoning about combinatorial interactions. We introduce MUGO (Multi-head Genomic Optimization), an in silico perturbation framework that casts variant discovery as differentiable combinatorial optimization over genomic sequence. MUGO relaxes discrete edits into a continuous probabilistic nucleotide mask and performs gradient-based optimization in input space to identify single- or multi-variant perturbations that maximize a user-specified molecular objective under a sequence-to-signal foundation model. This formulation makes genome-scale search computationally tractable while retaining direct, cell-type-specific molecular readouts and enabling precise quantification of non-additive interaction effects. Across five modalities and seven tissues using two foundation-model backbones, MUGO consistently outperforms three baselines in both optimization efficiency and effect modulation, while preserving robustness and cell-type specificity. Finally, MUGO-prioritized variants are enriched for GWAS signals across diverse tissues and a broad spectrum of complex traits, turning foundation-model predictions into scalable, cell-type-resolved hypotheses for causal variant discovery and combinatorial regulatory mechanisms.

*Corresponding author. This work was supported by the National Institutes of Health [R01HG012572, R01NS128523, R01DA063316].



This work is licensed under a Creative Commons Attribution 4.0 International License. *KDD '26, Jeju Island, Republic of Korea*
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2259-2/2026/08
<https://doi.org/10.1145/3770855.3818971>

CCS Concepts

• **Computing methodologies** → **Machine learning**; • **Applied computing** → **Computational biology**.

Keywords

Deep Learning, Combinatorial Optimization, Causal Variant Identification, Gradient Ascent, Discrete Signal Processing

ACM Reference Format:

Simon Dongbo Sun, Junhao Liu, Pengcheng Xu, Yaqi Hu, Martin Jinye Zhang, and Jing Zhang. 2026. MUGO: Differentiable Combinatorial Optimization for Causal Variant Discovery in the Non-coding Genome. In *Proceedings of the 32nd ACM SIGKDD Conference on Knowledge Discovery and Data Mining V.2 (KDD '26)*, August 09–13, 2026, Jeju Island, Republic of Korea. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3770855.3818971>

Resource Availability:

The source code for MUGO is publicly available at <https://github.com/aicb-ZhangLabs/MUGO> and is permanently archived on Zenodo at <https://doi.org/10.5281/zenodo.20740647>.

1 Introduction

Deciphering how genetic variants perturb gene regulation is central to translating human genetics into biological insight, improved diagnosis, and, ultimately, targeted therapies [11, 58]. Yet despite the success of genome-wide association studies (GWAS) in mapping thousands of disease-associated loci—most of which fall in the non-coding genome [7, 37]—these signals still seldom yield clear regulatory mechanisms, due to at least three persistent challenges. First, association peaks rarely yield cell-type-specific molecular interpretation, typically failing to pinpoint the perturbed regulatory element, the affected downstream program, or the direction and magnitude of effects across the specific cell types and states that mediate disease risk [8, 39]. Second, variant-to-gene causality remains difficult to establish, because linkage disequilibrium disperses statistical support across correlated variants and obscures

the true causal change and its target gene(s) [18, 26, 47]. Third, existing analyses typically consider variants in isolation, limiting combinatorial reasoning about how multiple variants jointly modulate shared regulatory circuitry [38, 52]. Overcoming these barriers will require unified frameworks that can resolve causal variants, attribute their cell-type-specific molecular consequences, and model higher-order regulatory interactions.

In response to these gaps, many computational tools have emerged to prioritize candidate variants and annotate their putative functions. Widely used predictors such as CADD and FunSeq integrate diverse genomic and regulatory features to score variant deleteriousness and highlight non-coding changes likely to perturb cis-regulatory programs [15, 29, 45], while conservation-based metrics such as GERP quantify evolutionary constraint to flag nucleotides under purifying selection [13, 50]. Resources such as KEGG provide pathway-level context for genes implicated by genetic studies [23]. In parallel, statistical fine-mapping approaches—including FINEMAP, CAVIAR, PAINTOR and SuSiE, among others—aim to narrow association signals into smaller credible sets by modeling linkage disequilibrium and, in some cases, incorporating functional priors [6, 18, 25, 26, 59]. Despite these advances, most approaches still lack cell-type-specific molecular interpretation, rarely resolve variant-to-gene causality, and cannot model combinatorial effects. Consequently, quantifying variant impact—and using it for reliable interpretation and prioritization—remains challenging [60, 61].

Recent progress in deep learning—particularly sequence-to-signal foundation models that map long genomic context to regulatory readouts—has broadened what is feasible for variant interpretation [14, 30]. Starting from early demonstrations that long-range DNA sequence can predict diverse genomic tracks (e.g., DeepSEA, Basenji, Enformer) [3, 24, 64], subsequent models have improved positional resolution and expanded the predicted modalities and cell contexts (e.g., Borzoi, Sei) [10, 34], with newer efforts moving toward unified, longer-context prediction across tissues and cell types at scale (e.g., Caduceus, Nucleotide Transformer, HyenaDNA, Evo) [12, 41, 42, 48]. By reducing reliance on hand-crafted features and instead learning regulatory patterns directly from sequence, these models make variant interpretation a concrete computational query: they support systematic *in silico* perturbation studies—from allele swaps and saturation mutagenesis to multi-variant edits—to quantify cell-type-specific molecular effects [57, 63], provide a basis for resolving variant-to-gene causality by tracing how allelic changes propagate through predicted regulatory intermediates to downstream gene regulation, and enable principled searches over combinatorial variant sets that amplify, buffer, or redirect regulatory programs [54]. Collectively, this paradigm offers a practical route to linking **sequence perturbation** with **molecular consequence** in a more direct and quantitative manner, thereby addressing the three challenges outlined above.

By enabling quantitative, cell-type-aware *in silico* sequence perturbations with direct readouts in predicted genomic tracks, sequence-to-signal foundation models make feature selection a natural next step: variants, motifs, or sequence segments can be edited and their molecular consequences assessed computationally. Permutation-based feature selection is therefore appealing for variant interpretation, but genome-scale applications introduce **three practical obstacles**. First, the search space is vast— $\sim 20,000$ genes $\times$ $\sim$

10^6 variants $\times$ multiple allele permutations—placing exhaustive perturbation beyond reach; genome-wide single-variant screens are already prohibitive, and higher-order combinatorial exploration is essentially impossible. Second, regulatory architecture is intrinsically combinatorial, so single-variant perturbations often miss synergy, buffering, and higher-order interactions [44, 62]. Third, many applications require global, context-robust attributions that capture patterns conserved across loci, cell types, and conditions, rather than explanations tailored to individual examples. Gradient-based shortcuts, including saliency maps, partially address computational cost but are limited by causal faithfulness, as they largely reflect local model sensitivity and can be unstable under saturation or small perturbations [28, 49, 51, 53]. These limitations motivate new methods that scale to genome-wide search while supporting combinatorial reasoning and more causally interpretable attributions.

To address these challenges, we introduce **MUGO (Multi-head Genomic Optimization)**, a framework that recasts causal variant discovery and prioritization as a **differentiable combinatorial optimization** problem. Rather than relying on discrete enumeration, MUGO uses a soft-masking relaxation inspired by recent progress in counterfactual explanation [27, 35]. Specifically, discrete DNA sequence edits are replaced with a continuous probabilistic mask over nucleotides, which makes the search space differentiable and enables end-to-end gradient ascent directly in input space [22, 36]. Candidate perturbations can then be iteratively refined toward a specified molecular objective, such as increasing or decreasing a target regulatory signal or gene expression program. We implement this strategy by coupling MUGO to the state-of-the-art open-source Borzoi architecture [34], enabling zero-shot identification of high-impact—potentially combinatorial—variant sets in a given molecular context. This design provides **three practical advantages**: (i) direct, assay-level readouts that support cell-type-specific interpretation beyond black-box scores; (ii) an explicit perturbation-based route to separate causal from merely associated variants; and (iii) a computationally efficient alternative to brute-force combinatorial search.

Our main contributions are summarized as follows:

- We propose MUGO, an efficient *in silico* perturbation-based framework that reframes variant prioritization as a controllable optimization problem, linking sequence edits to cell-type-specific molecular consequences and combinatorial (non-additive) interaction effects—thereby enabling mechanistic, experimentally actionable hypotheses beyond the existing static scoring schemes.
- Using two sequence-to-signal foundation models, we benchmark MUGO across five modalities and seven tissues, showing consistent gains over three baselines and robustness across backbones while preserving cell-type specificity and accurately quantifying combinatorial effects.
- To demonstrate utility for disease studies, we show that MUGO highlights disease-relevant variants across diverse tissues and a broad spectrum of complex traits, enabling context-specific mechanistic interpretation beyond proximity- or association-based prioritization.

2 Relation to Prior Work

Variant impact quantification or prioritization methods broadly fall into four lines: conservation or deleteriousness scores, functional impact models, population-based molecular associations, and statistical fine-mapping. We briefly review each and highlight shared limitations in cell-type-resolved interpretation, variant-to-gene causality, and combinatorial reasoning.

Conservation- and deleteriousness-based scores offer a scalable first pass for variant prioritization. Metrics such as phyloP, phastCons, and GERP, together with integrative predictors like CADD, distill evolutionary constraint and diverse genomic features into a single severity score [13, 29, 45, 50]. While useful for genome-wide triage, they are largely agnostic to cell type and state and rarely identify the perturbed regulatory process or target gene, limiting context-specific, testable hypotheses.

Functional impact models for coding and regulatory sequence aim to move beyond generic severity scores by estimating more specific consequences. For coding variants, tools such as SIFT, PolyPhen-2, and REVEL predict protein-level disruption [1, 21, 40], while non-coding methods including GWAVA, FunSeq2, and LIN-SIGHT leverage regulatory annotations and motif-perturbation signals [15, 20, 46]. However, many still depend on aggregated annotations or simplified sequence assumptions, limiting sensitivity to state-dependent regulation. Moreover, even when disruption is predicted, linking a variant to the causal element and downstream target gene often remains indirect [39].

Population-based molecular associations provide empirical support by linking genetic variation to molecular phenotypes. Resources such as GTEx map eQTL and sQTL effects across tissues [55, 56], and allele-specific analyses can further support regulatory hypotheses [9]. Colocalization methods (e.g., coloc, eCAVIAR) assess whether molecular and disease associations are consistent with a shared signal [17, 19]. However, limited cell-type and developmental resolution—and the lack of disease-relevant perturbation contexts—often weakens interpretation. Moreover, linkage disequilibrium (LD) blurs variant-level attribution, and shared signals alone do not establish a causal chain from variant to regulatory element to target gene [43].

Statistical fine-mapping refines association signals by modeling linkage disequilibrium to estimate posterior inclusion probabilities and credible sets (e.g., FINEMAP, SuSiE, CAVIAR, PAINTOR), sometimes aided by functional priors [6, 18, 25, 59]. Although it sharpens prioritization, credible sets often remain sizeable in highly correlated regions, and the resulting probabilities offer limited mechanistic insight [61]. Thus, fine-mapping narrows candidates but typically does not identify the perturbed regulatory program, affected gene, or context-dependent effects.

Perturbation-based Feature Selection and Model Interpretation With the rise of deep learning “oracles,” perturbation-based methods have become a standard tool for feature attribution and interpretation. *In silico mutagenesis* (ISM) [2, 64] estimates variant effects by mutating a reference base to an alternative allele and measuring the change in predicted output. However, ISM is computationally expensive $O(L)$ forward passes for a length- L sequence) and inherently local, evaluating one variant at a time

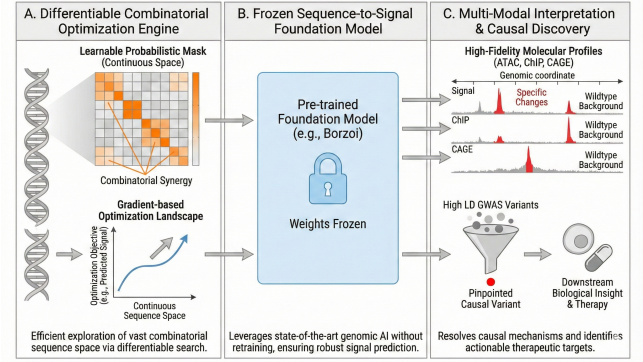


Figure 1: MUGO framework overview. (A) Engine: Gumbel-Softmax relaxation enables differentiable search over the combinatorial variant space. (B) Model: A straight-through estimator (STE) connects discrete DNA inputs to a frozen sequence-to-signal model while preserving gradient flow. (C) Discovery: Multi-modal objectives help disentangle causal variants from high-LD noise.

and thus missing combinatorial, context-dependent regulatory effects. Higher-order extensions quickly become intractable as the search space grows exponentially [30, 54]. To reduce the cost of explicit perturbation, gradient-based attribution methods such as Saliency Maps [51], DeepLIFT [49], and Integrated Gradients [53] backpropagate prediction signals to the input to estimate feature importance in $O(1)$ time. However, in genomics these attributions are often noisy, sensitive to baseline choices and saturation, and can suffer from “shattered gradients” in deep architectures such as Transformers [4, 16]. More fundamentally, gradients are local linear approximations and may mischaracterize the non-linear effects of discrete base substitutions. Here we retain the efficiency of gradient-based optimization while improving robustness and differentiability through soft-masking and consensus mechanisms, conceptually extending generative optimization approaches [27, 32, 35] to discrete variant prioritization.

3 Methodology

In this section, we present MUGO, a differentiable framework for causal variant discovery and objective-driven sequence perturbation. In contrast to perturbation-based strategies that rely on discrete enumeration, MUGO formulates the search as a variational inference problem, enabling efficient gradient-based optimization over high-dimensional combinatorial spaces. The method iteratively refines candidate perturbations toward a user-specified molecular objective, as summarized in Fig. 1.

3.1 Problem Formulation

3.1.1 Genomic Signal and Search Space. Let $\mathcal{S} = \{A, C, G, T\}$ denote the alphabet of nucleotides. We represent the reference genomic sequence as a one-hot encoded matrix $\mathbf{X}_{\text{ref}} \in \{0, 1\}^{L \times 4}$, where L is the sequence length (e.g., 524,288 bp for Borzoi). The goal of regulatory variant discovery is to identify specific mutations that significantly alter downstream biological signals, such as gene expression or chromatin accessibility.

We constrain our optimization to a biologically valid manifold. Let $\mathcal{V} = \{v_1, v_2, \dots, v_N\}$ be a set of candidate variants (e.g., common SNPs from the Geuvadis dataset), where each v_i is defined by a genomic position p_i and an alternative allele. This constraint is crucial to ensure that the identified variants are naturally occurring rather than synthetic adversarial perturbations.

3.1.2 Optimization Objective. Our objective is to learn a sparse binary mask $\mathbf{m} \in \{0, 1\}^N$ that selects a subset of variants from $\mathcal{V}$ to maximize a target function $\mathcal{F}$. Let $\mathbf{X}_{\text{alt}}$ denote the genomic sequence where all candidate variants in $\mathcal{V}$ are applied. The optimization problem can be formally stated as finding the optimal combination mask $\mathbf{m}^*$:

$$\mathbf{m}^* = \arg \max_{\mathbf{m} \in \{0, 1\}^N} \mathcal{J}(\mathbf{m}) = \mathcal{F}(\mathbf{X}_{\text{ref}} + \mathbf{m} \odot (\mathbf{X}_{\text{alt}} - \mathbf{X}_{\text{ref}})) \quad (1)$$

s.t. $\|\mathbf{m}\|_0 \leq k$

where $\mathcal{F}(\cdot)$ is a differentiable deep learning oracle (e.g., Borzoi) predicting gene expression, and k is a sparsity constraint controlling the number of selected variants. Directly solving Eq. (1) is an NP-hard combinatorial problem. The search space size is $\binom{N}{k}$, which grows exponentially with N . For a typical locus with $N = 1000$ variants, exhaustive search is computationally intractable.

3.2 Differentiable Variational Relaxation

To overcome the discreteness of the binary mask $\mathbf{m}$, we relax the problem into a continuous space using principles from **Variational Inference**. We introduce a learnable variational distribution $q_\phi(\mathbf{z})$ over the selection variables $\mathbf{z}$, parameterized by real-valued logits $\phi \in \mathbb{R}^N$.

3.2.1 Gumbel-Softmax Reparameterization. Standard categorical sampling $z \sim \text{Categorical}(\text{softmax}(\phi))$ involves a non-differentiable arg max operation, which blocks gradient flow during backpropagation. To enable end-to-end training, we employ the **Gumbel-Softmax relaxation** [22, 36]. This technique reparameterizes the sampling process by introducing independent Gumbel noise, allowing gradients to propagate through the stochastic samples.

Specifically, we generate a continuous relaxation vector $\mathbf{y} \in [0, 1]^N$ using the following transformation:

$$y_i = \frac{\exp((\phi_i + g_i)/\tau)}{\sum_{j=1}^N \exp((\phi_j + g_j)/\tau)} \quad (2)$$

where $g_i = -\log(-\log(u))$ with $u \sim \text{Uniform}(0, 1)$.

The temperature parameter $\tau > 0$ controls the entropy of the distribution. When $\tau \rightarrow \infty$, the samples approach a uniform distribution. As $\tau \rightarrow 0$, the distribution $q_\phi(\mathbf{y})$ converges to a concrete one-hot categorical distribution. In our framework, we employ an annealing schedule for τ , starting with a high temperature for exploration and gradually cooling it to encourage convergence to discrete solutions.

3.3 Bio-Plausible Gradient Estimation via STE

A unique challenge in genomic deep learning is the **Out-of-Distribution (OOD)** problem. Foundation models like Borzoi were trained strictly

on discrete DNA sequences (A, C, G, T). Feeding continuous, fractional inputs (e.g., a "soft" mixture of 0.6 A + 0.4 T) during optimization can lead to undefined behavior and artifacts in the model's internal activations.

To resolve this, we strictly enforce biological discreteness during the forward pass while maintaining differentiability during the backward pass. We utilize the **Straight-Through Estimator (STE)** [5].

Forward Pass (Discretization): We discretize the soft sample $\mathbf{y}$ into a hard binary mask $\mathbf{z}_{\text{hard}}$ using a greedy argmax or thresholding operation. This ensures that the input to the oracle $\mathcal{F}$ is always a valid, discrete genomic sequence:

$$\tilde{\mathbf{X}} = \mathbf{X}_{\text{ref}} + \mathbf{z}_{\text{hard}} \odot (\mathbf{X}_{\text{alt}} - \mathbf{X}_{\text{ref}}) \quad (3)$$

Backward Pass (Gradient Flow): Since the discretization step has zero derivative almost everywhere, we approximate the gradient of the hard sample with the gradient of the soft sample. This is implemented via the stop-gradient operator $\text{sg}(\cdot)$:

$$\mathbf{z}_{\text{input}} = \mathbf{z}_{\text{hard}} - \text{sg}(\mathbf{y}) + \mathbf{y} \quad (4)$$

In this formulation, $\mathbf{z}_{\text{input}} \approx \mathbf{z}_{\text{hard}}$ numerically, but $\nabla_\phi \mathbf{z}_{\text{input}} = \nabla_\phi \mathbf{y}$. This trick allows the gradients from the expression loss $\nabla \mathcal{J}$ to bypass the non-differentiable discretization step and update the variational parameters ϕ .

3.4 Stochastic Consensus for Variance Reduction

Optimization in discrete high-dimensional spaces is prone to high variance and local optima. A single stochastic sampling trajectory might be misled by gradient noise. To mitigate this, we propose a **Multi-Head Consensus Strategy**, which can be interpreted as a Monte Carlo variance reduction technique.

We maintain an ensemble of K independent optimization heads, each with its own noise trajectory. Instead of optimizing a single sample, we maximize the expected gain approximated by K Monte Carlo samples:

$$\nabla_\phi \mathbb{E}_q[\mathcal{J}] \approx \frac{1}{K} \sum_{k=1}^K \nabla_\phi \mathcal{F}_k(\mathbf{z}^{(k)}) \quad (5)$$

where $\mathbf{z}^{(k)}$ represents the variant mask sampled by the k -th head.

Mechanism. By averaging gradients across diverse heads, MUGO filters out stochastic noise (aleatoric uncertainty) inherent in the sampling process. True causal variants, which provide robust signal improvements, will receive consistent positive gradients across multiple heads, whereas spurious variants driven by random noise will have their gradients averaged out to zero. This averaging reduces stochastic gradient variance under independent sampling and stabilizes the learning process.

3.5 Optimization Algorithm and Complexity

The complete training procedure is outlined in Algorithm 1. We perform iterative gradient ascent on the variational parameters ϕ .

Complexity Analysis. Traditional exhaustive search requires $\mathcal{O}(\binom{N}{k})$ inferences. In contrast, MUGO requires $\mathcal{O}(T \times K)$ inferences, where T is the number of iterations (typically 100-200) and K is the ensemble size (e.g., 20). Since $T \times K \ll \binom{N}{k}$, our method achieves a significant speedup, enabling genome-wide application.

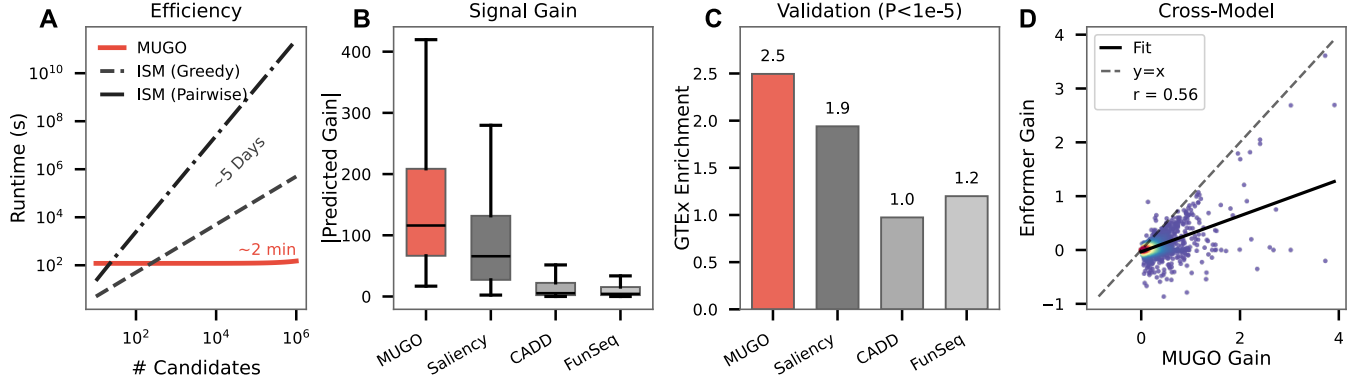


Figure 2: MUGO optimizes gene expression (summed RNA-track signal over exons) with superior efficiency and fidelity. (a) Scalability: MUGO achieves near-constant inference time, contrasting with the intractable scaling of linear (greedy) and combinatorial ISM strategies. **(b) Signal Gain:** MUGO consistently yields higher predicted gains compared to model-based baselines (Saliency) and non-model scoring methods (CADD, FunSeq). **(c) Validation:** Prioritized variants exhibit high accuracy in recovering biologically validated GTEx causal hits. **(d) Robustness:** Strong consistency ($r = 0.56$) between Borzoi and Enformer backbones indicates that MUGO captures intrinsic regulatory syntax rather than model-specific artifacts.

Algorithm 1 Multi-Head Genomic Optimization (MUGO)

Require: Reference Sequence $\mathbf{X}_{\text{ref}}$, Candidate Set $\mathcal{V}$, Oracle $\mathcal{F}$
Require: Hyperparams: Ensemble size K , Max steps T , Learning rate η

- 1: **Initialize** logits $\phi \in \mathbb{R}^N \leftarrow \mathbf{0}$, Temp $\tau \leftarrow \tau_{\max}$
- 2: Pre-compute difference tensor $\Delta\mathbf{X} = \mathbf{X}_{\text{alt}} - \mathbf{X}_{\text{ref}}$
- 3: **for** $t = 1$ to T **do**
- 4: Sample noise $\mathbf{g}^{(k)} \sim \text{Gumbel}(0, I)$ for $k = 1 \dots K$
- 5: **Relaxation:** $\mathbf{y}^{(k)} \leftarrow \text{Softmax}((\phi + \mathbf{g}^{(k)})/\tau)$
- 6: **STE Discretization:** $\mathbf{z}^{(k)} \leftarrow \text{Hard}(\mathbf{y}^{(k)}) - \text{sg}(\mathbf{y}^{(k)}) + \mathbf{y}^{(k)}$
- 7: **Forward Oracle:** $\mathcal{L} \leftarrow -\frac{1}{K} \sum_{k=1}^K \mathcal{F}(\mathbf{X}_{\text{ref}} + \mathbf{z}^{(k)} \odot \Delta\mathbf{X})$
- 8: **Backward:** Compute gradients $\nabla_{\phi} \mathcal{L}$
- 9: **Update:** $\phi \leftarrow \text{Adam}(\phi, \nabla_{\phi} \mathcal{L}, \eta)$
- 10: **Anneal:** $\tau \leftarrow \max(\tau \cdot \gamma, \tau_{\min})$
- 11: **end for**
- 12: **Return** Top- k variants indices from ϕ

3.6 Evaluation Framework

Regulatory variant discovery lacks a single definitive “ground truth.” To evaluate MUGO comprehensively, we adopt a multi-tiered framework spanning three key dimensions: Molecular Efficacy, Cross-Model Robustness, and Biological Causality.

3.6.1 Molecular Efficacy and Combinatorial Logic. To quantify the impact of optimized variants on molecular profiles, we define the **Predicted Gain** ($\Delta\mathbf{S}$) as the absolute difference in the oracle’s predicted signal (e.g., summed RNA-seq coverage) between the optimized sequence $\mathbf{X}_{\text{alt}}$ and the reference $\mathbf{X}_{\text{ref}}$.

To characterize non-linear interactions (Section 4.5), we define the **Interaction Ratio** (ρ). For a set of variants, ρ measures the ratio between the combined perturbation effect and the sum of individual marginal effects:

$$\rho = \frac{|\mathcal{F}(\mathbf{X}_{\text{comb}}) - \mathcal{F}(\mathbf{X}_{\text{ref}})|}{\sum_i |\mathcal{F}(\mathbf{X}_i) - \mathcal{F}(\mathbf{X}_{\text{ref}})|} \quad (6)$$

where $\rho \approx 1$ implies additivity, and $\rho > 1$ indicates synergistic (super-linear) logic.

3.6.2 Independent Validation using eQTL and GWAS. We validate the biological relevance of prioritized variants using two independent reference sets: expression quantitative trait loci (eQTLs) and disease-associated GWAS variants. eQTLs provide population-level evidence that variants modulate gene expression in vivo, while GWAS variants highlight loci implicated in disease and complex traits, allowing us to test whether prioritized perturbations concentrate in genetically supported regions.

Enrichment Score. We quantify performance using the **Enrichment Score** (ES), defined as the fold-increase in precision relative to random selection. Let $\hat{\mathcal{V}}_K$ be the top- K variants prioritized by the model, and $\mathcal{G}_{\text{truth}}$ be the ground-truth set (e.g., GTEx eQTLs or GWAS hits). The enrichment is calculated as:

$$ES@K = \frac{|\hat{\mathcal{V}}_K \cap \mathcal{G}_{\text{truth}}|/K}{|\mathcal{G}_{\text{truth}}|/N_{\text{total}}} \quad (7)$$

LD Resolution Stress Test. To evaluate causal discovery in high LD regions, we construct a “stress test” using fine-mapped GWAS data. For a known causal variant v_{causal} (PIP > 0.9), we identify a set of proxy variants $\mathcal{V}_{\text{proxy}}$ with high correlation ($r^2 > 0.8$). A valid interpretation method must assign a significantly higher importance score to v_{causal} than to any $v \in \mathcal{V}_{\text{proxy}}$, despite their statistical similarity.

4 Experiments

4.1 Experimental Setup

To systematically assess MUGO for decoding non-coding regulatory logic, we utilized state-of-the-art sequence-to-signal foundation models and large-scale human genetic datasets.

Datasets and Oracle Models. Our experiments focus on regulatory variation observed in human populations. We use the **Geuvadis** project [33] as the primary search space, restricting the candidate pool $\mathcal{V}$ to naturally occurring SNPs with minor allele frequency (MAF) > 0.05. For the predictive oracle, we primarily utilized **Borzoi** [34], leveraging its 524kb receptive field to capture long-range enhancer–promoter interactions. We conducted *in silico* perturbation studies across five modalities (including RNA-seq,

CAGE, ATAC-seq, ChIP-seq and DNase-seq) and seven tissues (e.g., Blood, Liver, Muscle). To ensure robustness, we replicated key results using **Enformer** [3], a Transformer-based model with distinct inductive biases, confirming that MUGO’s gradient-based search generalizes effectively to Transformer architectures.

Baselines. We benchmarked MUGO against three primary baselines used in the main comparisons, and included two auxiliary controls for scaling and stochastic-search behavior:

- **Perturbation-based:** *ISM* uses both Greedy (step-wise) and Exhaustive strategies. While considered as the “gold standard,” they suffer from prohibitive scaling ($O(N)$ or $O(N^2)$).
- **Gradient-based:** *Saliency Maps* [51]. Efficient but often noisy, approximating importance via input gradients ($|\nabla_x f(x)|$).
- **Stochastic Search:** *Random Search*, serving as a baseline to quantify the necessity of gradient guidance in sparse landscapes.
- **Scoring-based:** *CADD* and *FunSeq2*. Evolutionary and annotation-based scores widely used in clinical genetics but lacking cell-type context.

Evaluation Metrics. Following the framework defined in Section 3.6, we evaluated performance using:

- **Efficiency & Efficacy:** We report Wall-clock inference time and the *Predicted Gain* (ΔS) in gene expression.
- **Robustness:** We report the Pearson correlation (r) of optimization gains between Borzoi and Enformer architectures.
- **Ground Truth Validation:** We report the *Enrichment Score* (ES) against **GTEx v8 cis-eQTLs** ($P < 10^{-5}$) for molecular relevance, and against **GWAS Catalog** variants for disease relevance.
- **Fine-Mapping Resolution:** We compare the optimization scores of fine-mapped causal variants ($PIP > 0.9$) versus high-LD proxies ($r^2 > 0.8$) to assess de-noising capabilities.

Implementation and reproducibility details. All genomic coordinates were represented on the GRCh38 reference genome. For Borzoi, we extracted 524,288-bp input windows, and for Enformer we used the corresponding 196,608-bp input windows required by the model. Reference sequences were one-hot encoded over the four nucleotide channels. To keep the optimization biologically grounded, the candidate set was restricted to naturally occurring common SNPs from the Geuvadis / 1000 Genomes Phase 3 resource with minor allele frequency greater than 0.05 within each target regulatory window. Alternative alleles were applied at their native genomic positions, so each optimized perturbation corresponds to an observed human variant rather than a synthetic adversarial edit.

Objective construction. For gene-expression optimization, we used GENCODE v41 canonical transcript annotations to identify the 32-bp model output bins overlapping exonic regions of the target gene. The predicted gene-level signal was computed by summing the predicted coverage values across those exon-overlapping bins. This gene-centric aggregation prevents the objective from being dominated by nearby intergenic signal and keeps the optimization tied to the transcriptional output used in the eQTL validation. For other modalities, we selected tissue-matched model output tracks and evaluated the hard allele-applied sequences after optimization, rather than reporting scores from soft masks.

Table 1: Multi-modal Benchmark. Magnitude of Mean Gain ($|\Delta S| \pm \text{SEM}$). Bold indicates best performance.

Modality	Tissue	MUGO	Sal.	CADD	FunSeq
RNA-seq	Blood	174.8 $\pm$ 18.5	58.2 $\pm$ 33.4	39.7 $\pm$ 24.3	18.8 $\pm$ 12.0
	Brain	150.1 $\pm$ 13.3	19.8 $\pm$ 20.4	16.6 $\pm$ 11.9	2.1 $\pm$ 5.6
	Liver	162.2 $\pm$ 15.6	48.3 $\pm$ 20.5	14.3 $\pm$ 9.6	3.6 $\pm$ 4.0
	Heart	157.3 $\pm$ 13.7	37.1 $\pm$ 22.4	18.2 $\pm$ 12.8	4.9 $\pm$ 5.1
	Musc.	163.7 $\pm$ 12.8	43.4 $\pm$ 24.5	19.6 $\pm$ 14.0	0.4 $\pm$ 5.1
	Lung	108.8 $\pm$ 10.9	33.6 $\pm$ 16.2	6.0 $\pm$ 5.9	1.1 $\pm$ 4.2
	Kidney	130.6 $\pm$ 12.6	22.9 $\pm$ 15.1	11.4 $\pm$ 9.1	6.1 $\pm$ 3.9
CAGE	Blood	0.2 $\pm$ 0.0	0.0 $\pm$ 0.0	0.0 $\pm$ 0.0	0.0 $\pm$ 0.0
	Brain	3.4 $\pm$ 0.3	1.2 $\pm$ 0.7	0.6 $\pm$ 0.4	0.0 $\pm$ 0.2
	Liver	1.8 $\pm$ 0.2	0.5 $\pm$ 0.3	0.0 $\pm$ 0.1	0.0 $\pm$ 0.0
	Heart	1.1 $\pm$ 0.1	0.4 $\pm$ 0.3	0.2 $\pm$ 0.1	0.0 $\pm$ 0.0
	Musc.	2.3 $\pm$ 0.3	0.4 $\pm$ 0.4	0.2 $\pm$ 0.2	0.0 $\pm$ 0.1
ATAC	Blood	0.7 $\pm$ 0.1	0.1 $\pm$ 0.1	0.0 $\pm$ 0.0	0.0 $\pm$ 0.0
	Brain	0.8 $\pm$ 0.1	0.0 $\pm$ 0.1	0.0 $\pm$ 0.0	0.0 $\pm$ 0.0
DNase	Blood	5.2 $\pm$ 1.0	1.6 $\pm$ 1.0	0.2 $\pm$ 0.3	0.1 $\pm$ 0.2
	Brain	0.8 $\pm$ 0.2	0.1 $\pm$ 0.3	0.2 $\pm$ 0.2	0.0 $\pm$ 0.0
ChIP	Blood	26.6 $\pm$ 4.2	4.5 $\pm$ 4.9	0.5 $\pm$ 0.8	0.4 $\pm$ 1.0
	Brain	24.5 $\pm$ 3.2	4.5 $\pm$ 4.8	0.8 $\pm$ 1.0	1.0 $\pm$ 1.2

Validation protocol. For molecular validation, we used GTEx v8 cis-eQTLs with nominal $P < 10^{-5}$ as the external expression-supported variant set. Disease relevance was evaluated using GWAS Catalog variants in tissue-matched contexts. For the LD resolution stress test, fine-mapped variants with $PIP > 0.9$ were treated as causal candidates, and variants with $r^2 > 0.8$ to those candidates were used as high-LD proxies. This separates population-level enrichment from the more stringent LD de-noising task while keeping the validation criteria consistent with the metrics defined in Section 3.6.

Optimization settings. Unless otherwise specified, MUGO used Adam optimization with learning rate 1×10^{-2} , 200 optimization steps, an annealed Gumbel-Softmax temperature from 1.0 to 0.01, $K = 20$ optimization heads, top-5 variant selection, and gradient clipping with norm 1.0. These settings were selected to balance optimization stability and computational cost while preserving discrete, allele-valid forward passes through the frozen sequence-to-signal model.

4.2 MUGO Enables More Efficient Molecular Profile Modulation via Perturbations

We first evaluated MUGO’s capacity to modulate molecular profiles along **three complementary axes**: computational efficiency, perturbation efficacy, and robustness across predictive backbones.

Computational Efficiency. Distinct from standard feature-extraction methods, perturbation-based variant prioritization is fundamentally constrained by the size of the discrete search space: even under a conservative setting, it scales as $\sim 10^6$ variants $\times$ 3 candidate alleles $\times$ $\sim 20,000$ genes, with higher-order combinatorial searches expanding this space further. As a result, brute-force perturbation schemes such as ISM are prohibitive at scale—for example, ISM requires 5 days to prioritize a single variant for a single gene

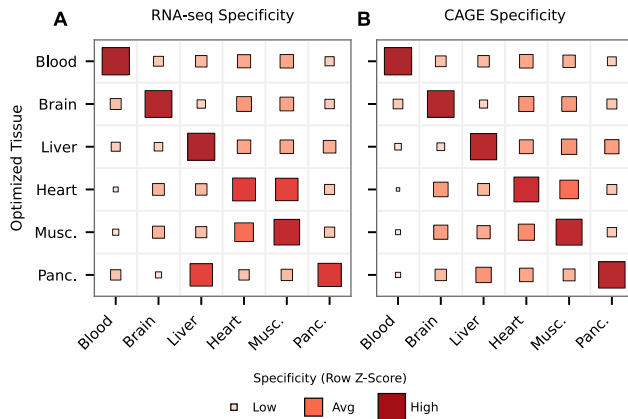


Figure 3: Multi-Modal Tissue Specificity. (a) RNA-seq and (b) CAGE evaluation matrices. The dominant diagonal signals confirm that MUGO-optimized variants yield maximal gains strictly within target tissues, demonstrating high cell-type specificity across expression modalities.

under a locus-level sweep, and an additional 500 days to identify a pair of variants that maximally modulates a target track (Fig. 2a). In contrast, MUGO replaces discrete enumeration with gradient ascent under a differentiable relaxation, enabling efficient exploration of combinatorial spaces that remain largely inaccessible to heuristic baselines. Consequently, MUGO achieves *near-constant inference time* (~ 2 minutes per locus, Fig. 2a) that is largely insensitive to the candidate set size, while preserving direct, model-based readouts through differentiable optimization. In practice, the shared-backbone implementation keeps the memory footprint manageable for $K = 20$ while scaling linearly in runtime.

Perturbation Efficacy. Given the prohibitive cost of brute-force perturbation, we benchmarked MUGO against representative methods from each major category of variant scoring and interpretation to assess efficacy in modulating predicted gene expression. As shown in Fig. 2b, approaches that explicitly extract actionable features for sequence editing consistently outperform non-interpretable prioritization scores in this setting. For example, MUGO and saliency-based selection achieved average gene expression changes of 174.80 and 58.20, respectively, after introducing selected mutations, exceeding model-agnostic scoring methods such as CADD and FunSeq (39.70 and 18.80; Fig. 2b). Moreover, within the feature-selection class, MUGO noticeably outperformed saliency-based baselines (174.80 vs. 58.20), highlighting the advantage of our differentiable optimization for driving targeted molecular profiles.

Because CADD and FunSeq do not operate on a sequence-to-signal model, we performed an additional validation using independent population-based molecular evidence to assess the physiological relevance of the prioritized variants. Specifically, we curated high-confidence cis-eQTLs from the GTEx v8 resource, which nominate variants associated with gene expression in large cohorts. If a method successfully identifies variants that modulate gene expression, its prioritized set should show increased overlap with these independently supported loci. We therefore quantified the enrichment of prioritized variants relative to random selection. As

shown in Fig. 2c, variants selected by MUGO achieved an enrichment of 2.50 over random, exceeding the enrichments observed for all baselines (1.94, 0.97, and 1.20 for saliency, CADD, and FunSeq, respectively). These results indicate that MUGO prioritizes variants with stronger population-level support for gene expression modulation.

Robustness across Predictive Backbones. Finally, to assess whether MUGO captures intrinsic regulatory grammar rather than model-specific artifacts of a single DNA-to-signal architecture, we performed cross-model validation. Specifically, we took variants optimized by MUGO using Borzoi and evaluated them in an independent model, Enformer, using the same loci and molecular objectives (see details in Section 3.6). We then compared the predicted gains produced by the two models to quantify cross-backbone concordance. As shown in Fig. 2d, predicted effects were strongly correlated across architectures ($r = 0.56$), supporting the view that the prioritized variants perturb regulatory patterns that are recognized consistently across distinct deep learning models.

4.3 MUGO Deciphers Regulatory Landscapes Across Modalities with Cell-Type Resolution

Regulatory programs are inherently multi-modal, with sequence variation propagating through coordinated changes in chromatin accessibility, histone modifications, and transcription. Moreover, these effects are often highly cell-type specific, reflecting differences in regulatory state and factor availability across cellular contexts. Motivated by these considerations, we extended our evaluation of MUGO to two additional axes: its ability to modulate molecular profiles consistently across distinct modalities, and its capacity to preserve cell-type-specific effects.

Ability for Multi-Modal Modulation We curated five modalities spanning key regulatory layers (RNA-seq, CAGE, ATAC-seq, DNase-seq, and ChIP-seq) to evaluate MUGO in a multi-modal setting (see Section 3.6). For each modality, we assessed each method’s ability to modulate targeted signal tracks across diverse tissues by introducing the corresponding sequence perturbations and measuring the resulting change in predicted signal tracks. As summarized in Table 1, model-based feature-selection approaches that leverage sequence-to-signal models consistently outperformed model-agnostic prioritization scores. In particular, MUGO and saliency-based baselines achieved average gains of 13.4- and 3.3-fold, respectively, relative to CADD and FunSeq. Notably, this improvement was not confined to a single assay or tissue: consistent gains were observed across **all 18 of 18** test scenarios. Furthermore, MUGO outperformed saliency-based selection in **18 out of 18** tasks, supporting the advantage of differentiable combinatorial optimization for robust multi-modal modulation. Together, these results indicate that MUGO generalizes across diverse regulatory readouts and maintains strong performance across tissue contexts, providing a practical route to targeted perturbation design beyond single-modality benchmarks.

Ability to Preserve Cell-Type Specificity We next asked whether MUGO preserves cell-type specificity, which is essential for mechanistic interpretation in regulatory genomics. Using two readouts—gene expression and CAGE-derived promoter activity—we optimized variants under a tissue-specific objective and then measured

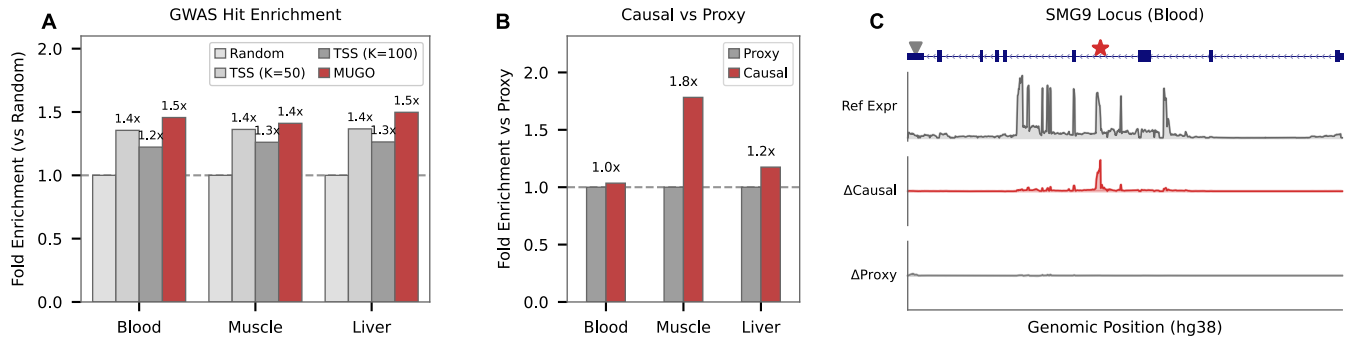


Figure 4: Resolving Variant-to-Gene Causality in High-LD Loci. (a) **GWAS Enrichment:** MUGO outperforms TSS-proximity and random baselines in recovering GWAS variants across tissues. (b) **LD Resolution:** In dense LD blocks, MUGO distinguishes causal variants (red) from correlated proxies (grey), effectively de-noising genetic signals. (c) **Case Study:** At the *SMG9* locus, the optimization signal (Δ Causal) selectively peaks at the causal variant, while the high-LD proxy (Δ Proxy) remains silent.

their predicted effects across a panel of other tissues. As shown in Fig. 3, responses were consistently strongest in the target tissue, producing a prominent diagonal in the cross-tissue effect matrix, whereas off-diagonal signals were attenuated. These results indicate that MUGO drives tissue-targeted molecular modulation with limited cross-context leakage, consistent with context-dependent regulatory grammar rather than broadly nonspecific sequence edits.

4.4 MUGO Establishes Direct Variant-to-Gene Links for Disease-Relevant Regulatory Programs

Interpreting disease-associated variation remains a central motivation for variant prioritization, yet GWAS loci often implicate broad LD blocks and rarely pinpoint the causal mutations or their mechanistic targets. By contrast, sequence-to-signal foundation models enable direct *in silico* perturbation of candidate alleles and quantitative readouts in predicted molecular tracks, providing a more direct route from sequence edits to molecular consequences. We therefore evaluated MUGO in a disease setting by testing its ability to identify key causal mutations across disease-relevant contexts.

GWAS Variants Enrichment. To evaluate the disease relevance of MUGO-prioritized variants, we measured their enrichment for tissue-matched GWAS variants relative to random expectation (Methods Section 3.6). We benchmarked against three baselines: random selection and two proximity heuristics that choose the top K variants closest to the transcription start site (TSS), with $K = 50$ and $K = 100$. Across all three tissues, MUGO consistently achieved the strongest enrichment—approximately 1.4–1.5 $\times$ over random—and outperformed other strategies (Fig. 4a). This consistent pattern indicates that MUGO does not simply favor variants near genes, but preferentially recovers disease-associated signals in the appropriate tissue context, supporting its use for disease-focused variant nomination.

De-noising Variants with High LD. We tested MUGO’s ability to distinguish causal variants from non-causal proxies in dense LD blocks ($r^2 > 0.8$), challenging the model to prioritize fine-mapped variants over highly correlated neighbors, which statistical association methods typically treat as nearly indistinguishable. In contrast to approaches that rely on population correlation or external annotations, MUGO scores variants using local sequence context (e.g.,

motif disruption). As a result, it assigns substantially higher optimization scores to causal candidates (red bars, 1.1 $\times$ to 1.8 $\times$) while down-weighting high-LD proxies (grey bars, Fig. 4b). This indicates that MUGO is sensitive to the specific disruption of regulatory grammar, rather than the variants’ correlation structure.

Case Studies of Causal Variants Identified by MUGO.

We next visualized this capability using track-level predictions at the *SMG9* locus (Figure 4c). In a standard GWAS association plot, this region would appear as a broad block of correlated signals. In contrast, MUGO yields a sharp and sparse optimization peak: the predicted signal change (Δ Causal) aligns with the fine-mapped variant, whereas a statistically linked proxy produces little to no response (Δ Proxy). This example highlights how deep learning-based optimization can break LD symmetry and nominate causal variants in settings where statistical genetics saturates.

4.5 MUGO Efficiently Uncovers Combinatorial Regulatory Logic

Finally, we explored non-linear variant-variant interactions within the regulatory code by probing for synergistic effects in transcriptional regulation. For example, transcription initiation often involves multiple transcription factors (TFs), and jointly perturbing these processes can produce effects that exceed the sum of individual disruptions. Such combinatorial effects are typically overlooked by existing baselines because the search space is prohibitively large. We therefore asked whether MUGO can fill in this gap to uncover this latent logic and exploit it to achieve larger optimization gains.

Identifying Synergy vs. Additivity. To quantify variant-variant interactions, we utilized the **Interaction Ratio** (ρ) (defined in Section 3.6), which normalizes the observed combinatorial gain by the sum of individual marginal gains.

Values $\rho > 1.0$ indicate positive synergy. In Figure 5a, we partitioned genes into a Synergistic tier (top-100 highest-gain genes) and a baseline tier, revealing a clear “synergy gap”: the Synergistic tier is shifted above 1.0 (mean $\rho > 1.2$, with a tail > 2.0), whereas the baseline tier is tightly centered near 1.0. These results indicate that, for the most responsive genes, MUGO identifies cooperative variant clusters rather than simply aggregating independently strong variants.

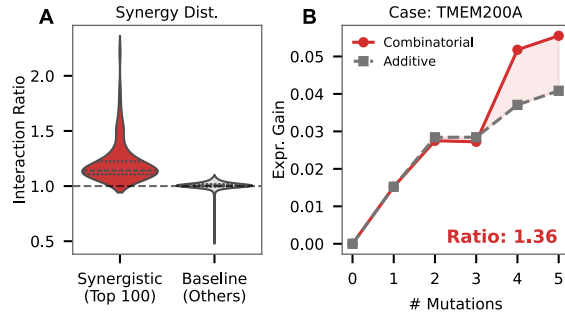


Figure 5: Uncovering Combinatorial Regulatory Logic. (a) Synergy: Top-100 synergistic genes (red) show interaction ratios > 1.0 , confirming cooperative effects beyond the additive baseline. (b) Trajectory: *TMEM200A* reveals super-linear gains (red) exceeding additive expectations (grey), implying threshold-based regulation.

To illustrate how such synergy emerges during optimization, we traced the stepwise trajectory for a representative high-synergy gene, *TMEM200A*. In Figure 5b, the grey dashed line denotes the expected additive gain (sum of individual effects), whereas the red curve shows the predicted response as variants are sequentially added.

At step 4, a specific variant triggers a super-linear spike exceeding additive expectations ($\rho = 1.36$, Fig. 5b). This suggests prioritized variants target cooperative motifs or redundant elements (e.g., shadow enhancers) [31]. Rather than functioning independently, these variants appear to engage cooperative mechanisms typically evolved to ensure phenotypic robustness and buffering against genetic perturbation; MUGO successfully identifies the specific combinatorial code required to overcome this buffering and drive a synergistic state shift.

5 Limitations and Ethical Considerations

MUGO depends on the fidelity and scope of the underlying genomic foundation models. As a result, its performance may vary in settings that are less well represented by the pretraining data or model architecture. In addition, because MUGO optimizes a model-defined molecular surrogate, the highest-scoring variants may not always translate to downstream phenotypes; its predictions remain hypothesis-generating and should ideally be followed by appropriate functional validation in relevant experimental systems. Increasing cell-state resolution, incorporating richer cellular context, and integrating improved sequence-to-track-to-phenotype models should further sharpen specificity and mechanistic interpretability.

More broadly, genomic AI models may reinforce existing biases when training and evaluation data do not adequately capture the diversity of human populations and biological contexts. Such limitations may arise from imbalances in ancestry, age, sex, gender-related factors, disease status, and environmental exposures. This study uses existing genomic resources and does not introduce new human-subject data collection, but downstream variant interpretation can still raise privacy, consent, and misuse concerns when applied to individual-level or clinically sensitive data. We therefore emphasize that MUGO should be applied with caution across diverse populations and settings, and that fairness, robustness, calibration, data

governance, and appropriate consent should be carefully assessed before any biomedical or clinical use.

6 Conclusion

In this work, we presented MUGO, a framework that reframes the challenge of causal variant discovery from a local attribution task to a global differentiable optimization problem. By bridging the discrete nature of genomic sequences with continuous gradient ascent via Gumbel-Softmax relaxation [22, 36] and Multi-Head Consensus, MUGO effectively navigates the combinatorial landscape of gene regulation that remains intractable for traditional perturbation baselines.

Our validation demonstrates significant advances in three key aspects. First, in terms of scalability, we achieve near-constant inference time, enabling genome-wide combinatorial screening orders of magnitude faster than *In Silico* Mutagenesis [2]. Second, regarding causality, we show that sequence-intrinsic optimization can disentangle physical drivers from statistical proxies in dense LD blocks, offering a powerful complement to population-based fine-mapping. Finally, we uncover combinatorial syntax, revealing that gene regulation is governed by non-linear synergistic interactions rather than additive effects.

As genomic foundation models continue to scale, the “signal optimization” paradigm introduced here offers a generalizable blueprint for decoding complex biological languages. Future work will extend this framework to design synthetic regulatory elements for gene therapy and explore broader multi-modal phenotypes beyond gene expression.

Reproducibility Statement

MUGO is intended to be reproducible from public resources and explicit optimization settings. All experiments use public genomic annotations, population variants, and regulatory resources, including GRCh38, Geuvadis / 1000 Genomes Phase 3 variants, GENCODE v41, GTEx v8 cis-eQTLs, GWAS Catalog variants, and public sequence-to-signal backbones. Candidate edits are restricted to observed common SNPs and evaluated as hard allele-applied sequences. The appendix summarizes the implementation choices, validation protocol, and hyperparameters needed to reproduce the camera-ready experiments. This keeps each reported result traceable to an explicit locus, candidate set, oracle track, and hard-selection rule. The same configuration also records the optimization seed and final top- k selection rule. Thus, rerunning a study on a new locus or modality preserves the same audit trail. Intermediate relaxed masks are not treated as final biological outputs. Only the final discrete allele set is used for reported validation and interpretation.

GenAI Disclosure

In the preparation of this work, we used Gemini to improve the readability and language quality of the manuscript. Additionally, Figure 1 was created with the assistance of Nano-banana to visually illustrate the conceptual framework. The authors reviewed and revised all AI-generated content and take full responsibility for the validity and accuracy of the final manuscript.

References

- [1] Ivan A Adzhubei, Steffen Schmidt, Leonid Peshkin, Vasily E Ramensky, Anna Gerasimova, Peer Bork, Alexey S Kondrashov, and Shamil R Sunyaev. 2010. A method and server for predicting damaging missense mutations. *Nature Methods* 7, 4 (2010), 248–249.
- [2] Babak Alipanahi, Andrew Delong, Matthew T. Weirauch, and Brendan J. Frey. 2015. Predicting the sequence specificities of DNA- and RNA-binding proteins by deep learning. *Nature Biotechnology* 33 (2015), 831–838.
- [3] Žiga Avsec, Vikram Agarwal, Daniel Visentin, et al. 2021. Effective gene expression prediction from sequence by integrating long-range interactions. *Nature Methods* 18 (2021), 1196–1203. doi:10.1038/s41592-021-01252-x
- [4] David Balduzzi et al. 2017. The Shattered Gradients Problem: If resnets are the answer, then what is the question? In *International Conference on Machine Learning (ICML)*. PMLR, Sydney, Australia, 342–350.
- [5] Yoshua Bengio, Nicholas Léonard, and Aaron Courville. 2013. Estimating or Propagating Gradients Through Stochastic Neurons for Conditional Computation. arXiv:1308.3432.
- [6] Christian Benner, Chris CA Spencer, Aki S Havulinna, Veikko Salomaa, Brian D Ripley, and Matti Pirinen. 2016. FINEMAP: efficient variable selection using summary data from genome-wide association studies. *Bioinformatics* 32, 10 (2016), 1493–1501.
- [7] Annalisa Buniello, Jacqueline AL MacArthur, Maria Cerezo, Laura W Harris, James Hayhurst, Cinzia Malangone, Aoife McMahon, Joannella Morales, Edward Mountjoy, Elliot Sollis, et al. 2019. The NHGRI-EBI GWAS Catalog of published genome-wide association studies, targeted arrays and summary statistics 2019. *Nucleic Acids Research* 47, D1 (2019), D1005–D1012.
- [8] Eddie Cano-Gamez and Gosia Trynka. 2020. From GWAS to function: using functional genomics to identify the mechanisms underlying complex human diseases. *Frontiers in Genetics* 11 (2020), 424.
- [9] Stephane E Castel, Ami Levy-Moonshine, Pejman Mohammadi, Eric Banks, and Tuuli Lappalainen. 2015. Tools and best practices for data processing in allelic expression analysis. *Genome Biology* 16 (2015), 1–12.
- [10] Kathleen M. Chen, Evan M. Cofer, Jian Zhou, et al. 2022. A sequence-based global map of regulatory activity for deciphering human genetics. *Nature Genetics* 54 (2022), 933–945. doi:10.1038/s41588-022-01102-2
- [11] Melina Claussnitzer, Judy H Cho, Rory Collins, Nancy J Cox, Mark J Daly, Manuel AR Ferreira, Daniel Godt, Eric S Lander, Daniel G MacArthur, Mark I McCarthy, et al. 2020. A brief history of human disease genetics. *Nature* 577, 7789 (2020), 179–189.
- [12] Hugo Dalla-Torre, Gonzalez Gonzalez, Javier Mendoza-Revilla, Nicolas L Caranza, Adam Henry Grzywaczewski, Francesco Oteri, Christian Dallago, Evan Trop, Hassan Sirelkhatim, et al. 2023. The Nucleotide Transformer: Building and Evaluating Robust Foundation Models for Human Genomics. bioRxiv preprint.
- [13] Eugene V Davydov, David L Goode, Marina Sirota, Gregory M Cooper, Arend Sidow, and Serafim Batzoglou. 2010. Identifying a high fraction of the human genome to be under selective constraint using GERP++. *PLoS Computational Biology* 6, 12 (2010), e1001025.
- [14] Gökçen Eraslan, Žiga Avsec, Julien Gagneur, and Fabian J Theis. 2019. Deep learning: new computational modelling techniques for genomics. *Nature Reviews Genetics* 20, 7 (2019), 389–403.
- [15] Yao Fu, Zhiping Liu, Shaokou Lou, Jason Bedford, Xinmeng Jasmine Mu, Kevin Y Yip, Ekta Khurana, and Mark Gerstein. 2014. FunSeq2: a framework for prioritizing noncoding regulatory variants in cancer. *Genome Biology* 15, 10 (2014), 480.
- [16] Amirata Ghorbani, Abubakar Abid, and James Zou. 2019. Interpretation of neural networks is fragile. *Proceedings of the AAAI Conference on Artificial Intelligence* 33, 01 (2019), 3681–3688.
- [17] Claudia Giambartolomei, Damjan Vukcevic, Eric E Schadt, Lude Franke, Aroon D Hingorani, Chris Wallace, and Vincent Plagnol. 2014. Bayesian test for colocalisation between pairs of genetic association studies using summary statistics. *PLoS Genetics* 10, 5 (2014), e1004383.
- [18] Farhad Hormozdiari, Emrah Kostem, Eun Yong Kang, Bogdan Pasaniuc, and Eleazar Eskin. 2014. Identifying causal variants at loci with multiple signals of association. *Genetics* 198, 2 (2014), 497–508.
- [19] Farhad Hormozdiari, Martijn Van De Bunt, Ayellet V Segrè, Xiao Li, Jong Wha J Joo, Alfred Bilow, Jae Hoon Sul, Sriram Sankararaman, Bogdan Pasaniuc, and Eleazar Eskin. 2016. Colocalization of GWAS and eQTL signals detects target genes. *American Journal of Human Genetics* 99, 6 (2016), 1245–1260.
- [20] Yi-Fei Huang, Brad Gulko, and Adam Siepel. 2017. Insight into the fitness cost of human mutation from molecular evolution. *Nature Genetics* 49, 6 (2017), 905–912.
- [21] Nilah M Ioannidis, Joseph H Rothstein, Vikas Pejaver, Sumit Middha, Shannon K McDonnell, Saurabh Baheti, Anthony Musolf, Qiongshi Li, Emily Holzinger, Deanna Karyadi, et al. 2016. REVEL: an ensemble method for predicting the pathogenicity of rare missense variants. *American Journal of Human Genetics* 99, 4 (2016), 877–885.
- [22] Eric Jang, Shixiang Gu, and Ben Poole. 2017. Categorical Reparameterization with Gumbel-Softmax. In *International Conference on Learning Representations (ICLR)*. OpenReview.net, Toulon, France, 13 pages.
- [23] Minoru Kanehisa, Miho Furumichi, Mao Tanabe, Yoko Sato, and Kanehisa Morishima. 2017. KEGG: new perspectives on genomes, pathways, diseases and drugs. *Nucleic Acids Research* 45, D1 (2017), D353–D361.
- [24] David R Kelley, Yakir A Reshef, and Maxwell Bileschi. 2018. Sequential regulatory activity prediction across chromosomes with convolutional neural networks. *Genome Research* 28, 5 (2018), 739–750.
- [25] Gleb Kichaev, Megan Roytman, Ruth Johnson, Eleazar Eskin, Sara Lindstrom, Peter Kraft, and Bogdan Pasaniuc. 2017. Improved methods for multi-trait fine mapping of pleiotropic risk loci. *Bioinformatics* 33, 2 (2017), 248–255.
- [26] Gleb Kichaev, Wan-Ping Yang, Sara Lindstrom, et al. 2014. Improving fine mapping of causal variants by integrating empirical candidates. *PLoS Genetics* 10, 10 (2014), e1004722.
- [27] Nathan Killoran, Leo J. Lee, Andrew Delong, David K. Duvenaud, and Brendan J. Frey. 2017. Generating and designing DNA with deep generative models. In *Advances in Neural Information Processing Systems (NeurIPS)*, Vol. 30. Curran Associates, Inc., Red Hook, NY, USA, 11 pages.
- [28] Pieter-Jan Kindermans, Kristof Schütt, Klaus-Robert Müller, and Sven Dähne. 2016. Investigating the influence of noise and distraction on neural network optimization. arXiv:1611.04552.
- [29] Martin Kircher, Daniela M Witten, Preti Jain, Brian J O’Roak, Gregory M Cooper, and Jay Shendure. 2014. A general framework for estimating the relative pathogenicity of human genetic variants. *Nature Genetics* 46, 3 (2014), 310–315.
- [30] Peter K Koo and Matt Ploenzke. 2021. Representation learning of genomic sequence motifs with convolutional neural networks. *PLoS Computational Biology* 17, 7 (2021), e1009131.
- [31] Evgeny Z. Kvon, Rachel Waymack, Mario Gad, and Zeba Wunderlich. 2021. Enhancer redundancy in development and disease. *Nature Reviews Genetics* 22 (2021), 324–336. doi:10.1038/s41576-020-00311-x
- [32] Jack Lanchantin, Ritambhara Singh, Beilun Wang, and Yanjun Qi. 2017. Deep Motif Dashboard: Visualizing and Understanding Genomic Sequences Using Deep Neural Networks. *Pacific Symposium on Biocomputing* 22 (2017), 254–265.
- [33] Tuuli Lappalainen, Michael Sammeth, Marc R Friedländer, Peter AC ’t Hoen, Jean Monlong, Manuel A Rivas, Mar González-Porta, Scott Kuersten, Thasso Griebel, Pedro G Ferreira, et al. 2013. Transcriptome and genome sequencing uncovers functional variation in humans. *Nature* 501, 7468 (2013), 506–511.
- [34] J. Linder, D. Srivastava, H. Yuan, et al. 2025. Predicting RNA-seq coverage from DNA sequence as a unifying model of gene regulation. *Nature Genetics* 57 (2025), 949–961. doi:10.1038/s41588-024-02053-6
- [35] Junhao Liu, Pengpeng Zhang, Martin Renqiang Min, and Jing Zhang. 2025. Identifying Combinatorial Regulatory Genes for Cell Fate Decision via Reparameterizable Subset Explanations. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD)*. Association for Computing Machinery, New York, NY, USA, 1823–1832. doi:10.1145/3711896.3737000
- [36] Chris J. Maddison, Andriy Mnih, and Yee Whye Teh. 2017. The Concrete Distribution: A Continuous Relaxation of Discrete Random Variables. In *International Conference on Learning Representations (ICLR)*. OpenReview.net, Toulon, France, 12 pages.
- [37] Matthew T Maurano, Richard Humbert, Eric Rynes, Robert E Thurman, Eric Haugen, Hao Wang, Alex P Reynolds, Richard Sandstrom, Hongzhu Qu, Jennifer Brody, et al. 2012. Systematic localization of common disease-associated variation in regulatory DNA. *Science* 337, 6099 (2012), 1190–1195.
- [38] Ekaterina Morgunova and Jussi Taipale. 2015. Structural insights into the cooperative binding of transcription factors to DNA. *Current Opinion in Structural Biology* 35 (2015), 1–9.
- [39] Joseph Nasser, Drew T Bergman, Charles P Fulco, Philine Guckelberger, Benjamin R Doughty, Tejal A Patwardhan, Thouis R Jones, Tung H Nguyen, Jacob C Ulirsch, Fritz Leekschas, et al. 2021. Genome-wide enhancer maps link risk variants to disease genes. *Nature* 593, 7858 (2021), 238–243.
- [40] Pauline C Ng and Steven Henikoff. 2003. SIFT: predicting amino acid changes that affect protein function. *Nucleic Acids Research* 31, 13 (2003), 3812–3814.
- [41] Eric Nguyen, Michael Poli, Matthew Durrant, et al. 2024. Sequence modeling and design from molecular to genome scale with Evo. *Science* 386 (2024), eado9336.
- [42] Eric Nguyen, Michael Poli, Marjan Faiz, Armin Thomas, et al. 2023. HyenaDNA: Long-Range Genomic Sequence Modeling at Single Nucleotide Resolution. In *Advances in Neural Information Processing Systems (NeurIPS)*, Vol. 36. Curran Associates, Inc., Red Hook, NY, USA, 18 pages.
- [43] Bogdan Pasaniuc and Alkes L. Price. 2017. Dissecting the genetics of complex traits using summary association statistics. *Nature Reviews Genetics* 18 (2017), 117–127.
- [44] Franziska Reiter, Sebastian Wienerroither, and Alexander Stark. 2017. Combinatorial function of transcription factors and cofactors. *Current Opinion in Genetics & Development* 43 (2017), 39–46.
- [45] Philipp Rentzsch, Daniela Witten, Gregory M Cooper, Jay Shendure, and Martin Kircher. 2019. CADD: predicting the deleteriousness of variants throughout the human genome. *Nucleic Acids Research* 47, D1 (2019), D886–D894.
- [46] Graham RS Ritchie, Ian Dunham, Eleftheria Zeggini, and Paul Flicek. 2014. Functional annotation of noncoding sequence variants. *Nature Methods* 11, 3 (2014),

- 294–296.
- [47] Daniel J Schaid, Wei Chen, and Nicholas B Larson. 2018. Statistical methods for identifying causal variants of complex traits. *Annual Review of Genomics and Human Genetics* 19 (2018), 391–419.
 - [48] Yair Schiff, Chia-Hsiang Kao, Aaron Gokaslan, et al. 2024. Caduceus: Bi-directional Equivariant Long-Range DNA Modeling. *International Conference on Machine Learning (ICML)* 235 (2024), 17 pages.
 - [49] Avanti Shrikumar, Peyton Greenside, and Anshul Kundaje. 2017. Learning Important Features Through Propagating Activation Differences. In *Proceedings of the 34th International Conference on Machine Learning (ICML)*. PMLR, Sydney, Australia, 3145–3153.
 - [50] Adam Siepel, Gill Bejerano, Jakob S Pedersen, Angie S Hinrichs, Minmei Hou, Kate Rosenbloom, Hiram Clawson, John Spieth, LaDeana W Hillier, Stephen Richards, et al. 2005. Evolutionarily conserved elements in vertebrate, insect, worm, and yeast genomes. *Genome Research* 15, 8 (2005), 1034–1050.
 - [51] Karen Simonyan, Andrea Vedaldi, and Andrew Zisserman. 2013. Deep Inside Convolutional Networks: Visualising Image Classification Models and Saliency Maps. arXiv:1312.6034.
 - [52] François Spitz and Eileen EM Furlong. 2012. Transcription factors: from enhancer binding to developmental control. *Nature Reviews Genetics* 13, 9 (2012), 613–626.
 - [53] Mukund Sundararajan, Ankur Taly, and Qiqi Yan. 2017. Axiomatic Attribution for Deep Networks. In *Proceedings of the 34th International Conference on Machine Learning (ICML)*. PMLR, Sydney, Australia, 3319–3328.
 - [54] Il Taskiran, K Spanier, et al. 2021. Cell-type-specific cis-regulatory elements in genetic disease constructs. *Nature Genetics* 53 (2021), 10 pages.
 - [55] The GTEx Consortium. 2013. The Genotype-Tissue Expression (GTEx) project. *Nature Genetics* 45 (2013), 580–585.
 - [56] The GTEx Consortium. 2020. The GTEx Consortium atlas of genetic regulatory effects across human tissues. *Science* 369, 6509 (2020), 1318–1330.
 - [57] Shushan Toneyan, Ziqi Tang, and Peter K Koo. 2022. Evaluating deep learning for predicting epigenomic profiles. *Nature Machine Intelligence* 4 (2022), 766–773.
 - [58] Peter M Visscher, Naomi R Wray, Qian Zhang, Pamela Sklar, Mark I McCarthy, Matthew A Brown, and Jian Yang. 2017. 10 years of GWAS discovery: biology, function, and translation. *American Journal of Human Genetics* 101, 1 (2017), 5–22.
 - [59] Gao Wang, Abhishek Sarkar, Peter Carbonetto, and Matthew Stephens. 2020. A simple new approach to variable selection in regression, with application to genetic fine-mapping. *Journal of the Royal Statistical Society Series B: Statistical Methodology* 82, 5 (2020), 1273–1300.
 - [60] Christopher Weeks, Henry Lu, Chen Wang, Ruoyao Sun, Shushan Toneyan, Gökcen Eraslan, Žiga Avsec, et al. 2023. Leveraging deep learning to understand the genetics of human disease. *Nature Reviews Genetics* 24 (2023), 1–21.
 - [61] Omer Weissbrod, Farhad Hormozdiari, Christian Benner, Ran Cui, Jacob Ulirsch, Steven Gazal, Armin P Schoech, Bryce Van De Geijn, Yakir Reshef, Carla Márquez-Luna, et al. 2020. Functionally informed fine-mapping and polygenic localization of complex trait heritability. *Nature Genetics* 52, 12 (2020), 1355–1363.
 - [62] Julia Zeitlinger. 2023. Combinatorial transcription factor binding and the problem of non-coding variant interpretation. *Current Opinion in Genetics & Development* 80 (2023), 102047.
 - [63] Anna Zhivi, Xiaoting Zhong, et al. 2023. Genome-wide prediction of disease variants with a deep learning framework. *Nature Communications* 14 (2023), 1–12.
 - [64] Jian Zhou and Olga G Troyanskaya. 2015. Predicting effects of noncoding variants with deep learning-based sequence model. *Nature Methods* 12, 10 (2015), 931–934.

A Additional Reproducibility Details

Data and model inputs. All experiments used GRCh38 coordinates. Borzoi inputs were 524,288 bp and Enformer inputs were 196,608 bp. Candidate variants were common SNPs from Geuvadis 1000 Genomes Phase 3 with MAF > 0.05 within the target regulatory windows. Reference sequences were one-hot encoded over the four nucleotide channels and centered on the target regulatory interval required by the corresponding oracle model. For each candidate SNP, the alternative allele was inserted at its native genomic position, preserving the surrounding haplotype context used by the sequence-to-signal backbone.

Gene-level aggregation. For each target gene, canonical transcript annotations from GENCODE v41 were used to identify exon-overlapping 32-bp output bins. Gene-level predicted expression was

computed as the sum of model-predicted coverage over those bins: $y_{\text{gene}} = \sum_{b \in \mathcal{B}_{\text{exons}}} \hat{y}_b$.

Track-level objectives. For non-expression modalities, the target score was computed from tissue-matched model output tracks. The optimization objective was always defined on a frozen oracle model, and model parameters were not updated during MUGO optimization. This design separates the search procedure from model training: MUGO changes only the candidate allele mask, while the underlying sequence-to-signal function remains fixed. After optimization, final scores were recomputed using hard allele-applied sequences, not the continuous relaxation used to obtain gradients.

Objective direction. The default experiments optimize for increased predicted signal in the target track, but the same formulation can be used for repression by reversing the objective sign. In either case, the candidate set, oracle model, and final hard-evaluation rule remain unchanged. This makes the optimization direction an explicit experimental choice rather than a property of the implementation. For all reported camera-ready experiments, the objective direction was matched to the molecular gain or enrichment metric described in the main text.

Tissue matching. Tissue-specific analyses used output tracks aligned to the target context whenever such tracks were available in the oracle model. The same selected alleles were then evaluated across other tissues for specificity analyses. This cross-tissue evaluation uses the model outputs only as readouts after selection, so it does not give the optimizer access to the off-target tissues when the target-tissue objective is being optimized. This distinction is important for interpreting the diagonal patterns in the tissue-specific evaluation matrices.

Discrete relaxation and hard evaluation. MUGO uses a Gumbel-Softmax relaxation to make the variant-selection mask differentiable during optimization. At each step, the relaxed mask provides gradients for updating the variant logits, while the forward sequence is discretized so that the oracle receives a valid nucleotide sequence. The final reported perturbation set is obtained by selecting the top-ranked variants under the learned logits and applying their alternative alleles directly. This prevents the evaluation from depending on fractional nucleotide states and keeps the reported effects comparable to standard in silico mutagenesis.

Multi-head consensus. The multi-head strategy runs several independent stochastic relaxations in parallel and averages their optimization signal. Each head samples its own Gumbel noise trajectory, which reduces the chance that a single noisy relaxation dominates the selected variant set. Variants that consistently improve the objective across heads receive reinforced gradients, whereas variants selected only by noise tend to be averaged down. In practice, this consensus mechanism improves stability without changing the frozen oracle or the biological interpretation of the selected alleles.

Baseline alignment. All baselines were evaluated on the same target loci and candidate variant pools used by MUGO. Perturbation baselines score variants by directly applying sequence edits and measuring the corresponding oracle response. Gradient baselines use local sensitivity information around the reference sequence, and annotation-based baselines rank candidates without access

to the tissue-specific oracle objective. This shared candidate-pool design keeps the comparison focused on the prioritization strategy rather than differences in variant availability.

Validation. GTEx v8 cis-eQTLs with nominal $P < 10^{-5}$ were used for molecular enrichment validation. GWAS Catalog variants were used for disease enrichment. For LD resolution, fine-mapped variants with PIP > 0.9 were compared against high-LD proxies with $r^2 > 0.8$.

LD-aware evaluation. The LD resolution analysis was designed to distinguish sequence-level functional evidence from population correlation. Fine-mapped variants with high posterior inclusion probability were treated as candidate causal variants, while highly correlated neighbors were used as proxy variants. MUGO and the baselines were then evaluated by comparing their assigned perturbation scores for the causal candidates versus the high-LD proxies. This protocol does not require new experimental labels beyond the fine-mapping and LD resources; instead, it asks whether the sequence-to-signal objective can prioritize the variant whose local sequence change is most compatible with the predicted molecular effect.

Reporting conventions. Runtime comparisons were reported at the locus level because perturbation methods differ sharply in how they scale with candidate set size. Enrichment scores were computed relative to random selection under the same candidate pool, so they should be interpreted as prioritization enrichment rather than absolute causal probability. Cross-backbone robustness was assessed by optimizing with one oracle and evaluating the selected alleles with an independent architecture under matched loci and molecular objectives. This reporting convention keeps the optimization procedure, validation resource, and biological interpretation separated.

Numerical stability. The temperature schedule starts with a smoother relaxation and gradually moves toward a sharper mask. This improves early exploration while still producing discrete top-ranked variants by the end of optimization. Gradient clipping is used to avoid occasional large updates from dominating the logits, especially when the oracle response changes sharply near a regulatory motif. These stability choices affect the optimization trajectory but not the final rule for constructing the evaluated sequence.

Implementation notes. The implementation stores reference and alternative alleles as paired tensors, so applying a variant corresponds to replacing the reference channel with the alternative channel at the corresponding genomic coordinate. Gradient updates are applied only to the variant logits. The sequence tensor passed to the oracle is reconstructed from the reference sequence and the selected allele mask at each optimization step. For memory efficiency, multiple heads share the same frozen oracle forward pass structure while maintaining independent mask samples.

Resource accounting. The dominant computational cost is the frozen oracle forward and backward pass through long genomic sequence windows. Because MUGO optimizes logits over a candidate set rather than enumerating every subset, increasing the number of candidate variants affects the mask dimension but does not require

Table A1: Optimization hyperparameters used in the main experiments.

Parameter	Value
Optimizer	Adam
Learning rate	1×10^{-2}
Batch size	1
Max iterations	200
Initial temperature	1.0
Minimum temperature	0.01
Annealing schedule	Exponential decay
Ensemble heads (K)	20
Sparsity constraint	Top-5 variants
Gradient clipping	1.0

a separate oracle call for every possible edit combination. The number of optimization steps and heads therefore provides the primary practical control over runtime and memory use.

Randomness and repeatability. Stochasticity enters through the Gumbel samples used by the relaxation and through any randomized baseline sampling procedure. For repeatable runs, the random seed, candidate variant order, temperature schedule, number of heads, and top- k selection rule should be recorded together. When comparing methods, random baselines should be resampled under the same candidate-pool constraints used by the optimized methods.

Output interpretation. MUGO returns a prioritized set of candidate alleles and an oracle-predicted molecular effect. These outputs should be interpreted as mechanistic hypotheses generated by a sequence-to-signal model, not as direct clinical assertions. Downstream use should therefore preserve the distinction between model-predicted molecular modulation, population-level enrichment, and experimentally validated function. This distinction is especially important in disease contexts where LD structure and cohort composition can affect external validation resources.

Reproducibility checklist. To reproduce a run, one must specify the genome build, target locus, candidate SNP list, oracle model, molecular track, target gene or output interval, optimization steps, temperature schedule, number of heads, sparsity level, and final hard-selection rule. The main experiments use the settings summarized in Table A1. When applying MUGO to a new tissue or molecular modality, the key change is the chosen oracle output track; the discrete optimization procedure and final hard-allele evaluation remain unchanged.

Artifact organization. The code release is organized around reusable configuration files rather than manuscript-specific scripts. Each run records the target interval, oracle backbone, selected output track, candidate set, and optimization hyperparameters, allowing the same workflow to be rerun on a different locus or modality with minimal changes. This structure is intended to make camera-ready results traceable from a concise experimental configuration to the final ranked allele set.