



Transformers interpreted as descent algorithms

- Algorithmic unrolling interprets transformer layers as optimization steps, therefore loss should decrease layerwise.
- Standard ERM training:

$$\mathbf{T}_U^* = \operatorname{argmin}_{\mathbf{T}} \mathbb{E} [f(\mathbf{X}, \Phi(\mathbf{X}; \mathbf{T}))] \quad (\text{ERM})$$

- Problem:** ERM training optimizes only for last-layer value of loss function f . The resulting **unconstrained transformer** is non-monotonic across layers.
- Leads to decreased robustness to OOD perturbations.

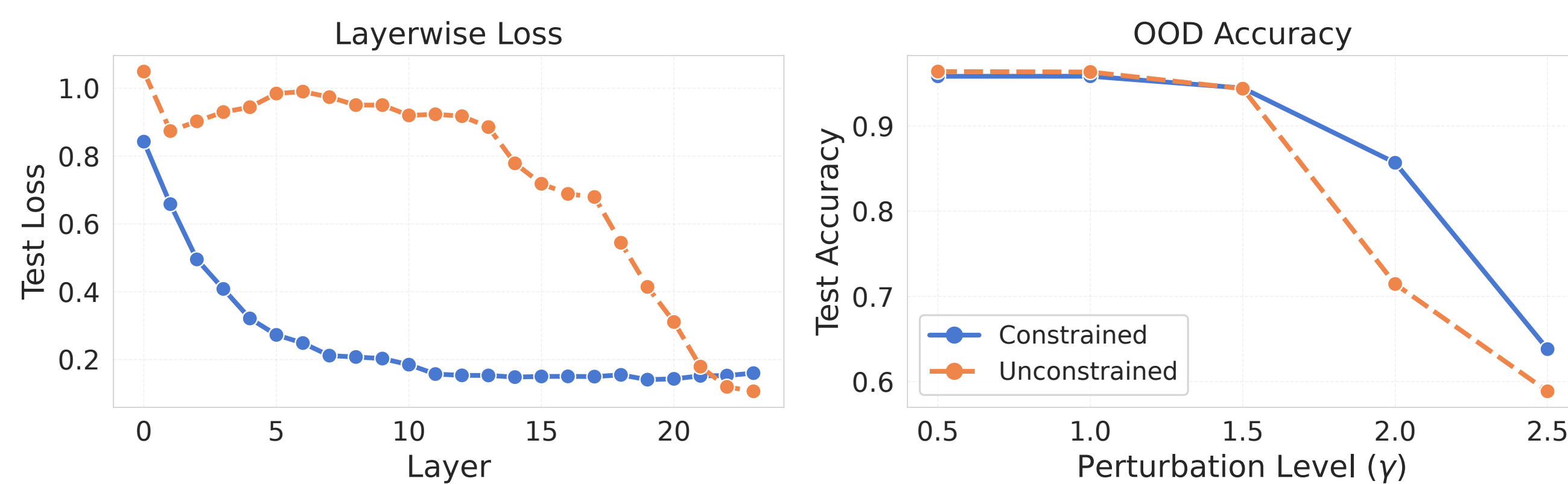


Figure 1. **Unconstrained** transformers show non-monotonic loss, **constrained** models descend smoothly with better OOD robustness. (Left: $\downarrow$, Right: $\uparrow$)

Training transformers to descend improves OOD robustness

- Key Idea:** We impose constraints to decrease loss monotonically in expectation, inspired by descent algorithms.
- Consider transformer given by layers

$$\mathbf{Z}_l = \mathbf{V}_l \mathbf{X}_{l-1} \times \operatorname{softmax}[(\mathbf{Q}_l \mathbf{X}_{l-1})^\top (\mathbf{K}_l \mathbf{X}_{l-1})]$$

$$\Phi_l(\mathbf{X}; \mathbf{T}) = \sigma[\mathbf{W}_l \mathbf{Z}_l + \mathbf{U}_l \mathbf{X}_{l-1}]$$

- Train each layer Φ_l to decrease f with step size $0 < \alpha < 1$:

Constrained Unrolled Transformer

$$\begin{aligned} \mathbf{T}^* &= \operatorname{argmin}_{\mathbf{T}} \mathbb{E} [f(\mathbf{X}, \Phi(\mathbf{X}; \mathbf{T}))] \\ \text{s.t. } \mathbb{E} [f(\mathbf{X}, \Phi_l(\mathbf{X}; \mathbf{T}))] &\leq (1 - \alpha_l) \mathbb{E} [f(\mathbf{X}, \Phi_{l-1}(\mathbf{X}; \mathbf{T}))], \forall l \end{aligned} \quad (\text{P-CUT})$$

- Lagrangian of (P-CUT):

$$\mathcal{L}(\mathbf{T}, \boldsymbol{\lambda}) = \mathbb{E} [f(\mathbf{X}, \Phi(\mathbf{X}; \mathbf{T}))] + \sum_{l=1}^L \lambda_l \mathbb{E} [f(\mathbf{X}, \Phi_l(\mathbf{X}; \mathbf{T})) - (1 - \alpha_l) f(\mathbf{X}, \Phi_{l-1}(\mathbf{X}; \mathbf{T}))],$$

- Optimize (P-CUT) by alternating primal and dual steps on the empirical dual problem:

$$\hat{D}^* = \max_{\boldsymbol{\lambda} \geq 0} \min_{\mathbf{T}} \hat{\mathcal{L}}(\mathbf{T}, \boldsymbol{\lambda})$$

Convergence & OOD Generalization guarantees

We bound the optimality gap $\Delta_k^* := f(\mathbf{X}, \Phi_k(\mathbf{X}; \mathbf{T}^*)) - f(\mathbf{X}, \mathbf{Y}^*)$.

- Theorem 1 (Convergence):** Constrained transformers attain near-optimal loss:

$$\lim_{l \rightarrow \infty} \min_{k \leq l} \mathbb{E} [\Delta_k^*] \leq \frac{1}{\alpha} \left(\zeta(M, \delta) + \frac{C\delta\nu}{1 - \delta} \right)$$

- Theorem 2 (OOD Generalization):** Under shifted distribution $D_{x'}$

$$\lim_{l \rightarrow \infty} \min_{k \leq l} \mathbb{E}_{D_{x'}} [\Delta_k^*] \leq \frac{1}{\alpha} \left(\zeta(M, \delta) + \frac{C\delta\nu}{1 - \delta} + C\tau \right)$$

where $\tau = d(D_x, D_{x'}) + d(D_{x'}, D_x)$ measures distribution shift.

- Bounds depend on sample complexity, expressivity, distribution distance and step size α .

Unrolled text classifiers are robust to noisy embeddings

- Leverage robustness to **distribution shifts** to classify with perturbed embeddings $\tilde{\mathbf{X}}$.
- Evaluate layerwise loss by attaching a shared readout layer ψ .

$$f(\tilde{\mathbf{X}}, \mathbf{q}, \Phi(\tilde{\mathbf{X}}; \mathbf{T})) = - \sum_{c=1}^C q_c \log \psi(\Phi(\tilde{\mathbf{X}}; \mathbf{T}))_c$$

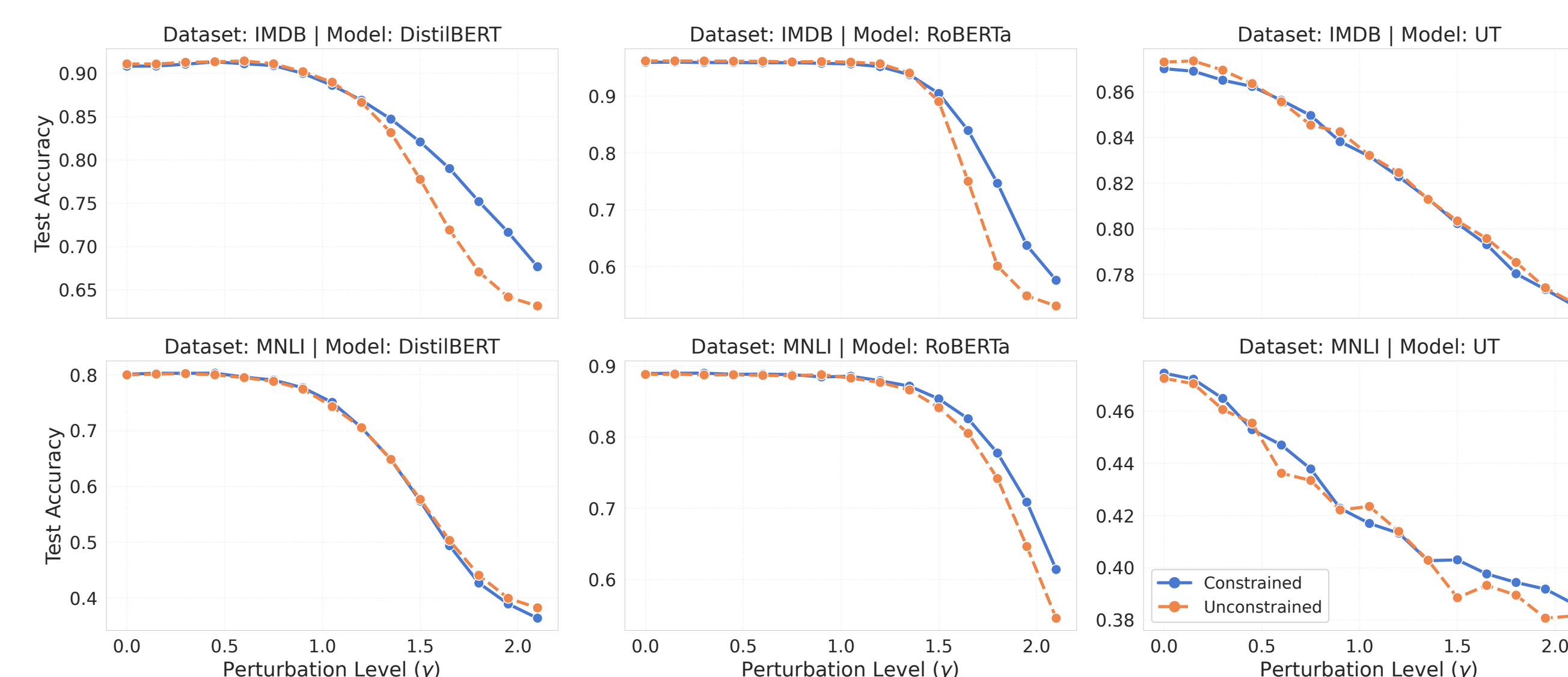


Figure 2. Accuracy vs. test perturbation γ ($\uparrow$ higher is better).

- Similar results with uniform noise and discrete perturbations:

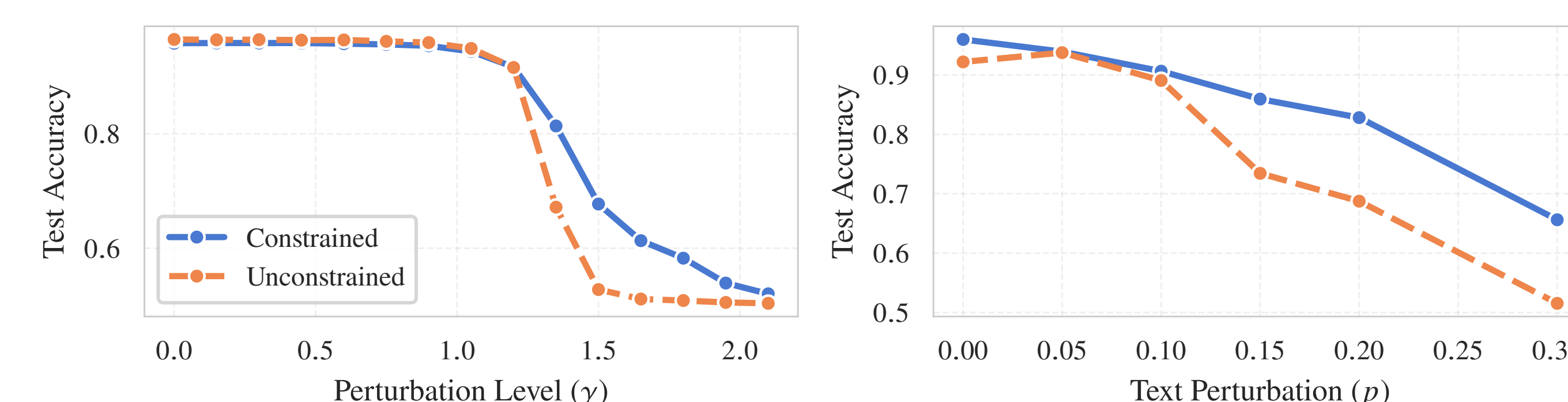


Figure 3. Uniformly distributed perturbations (left) and random text perturbations (character and word level, right).

LLM supervised fine-tuning: Improves OOD, Maintains ID

- Llama 3.1 8B SFT on Alpaca dataset with perturbed token embeddings.

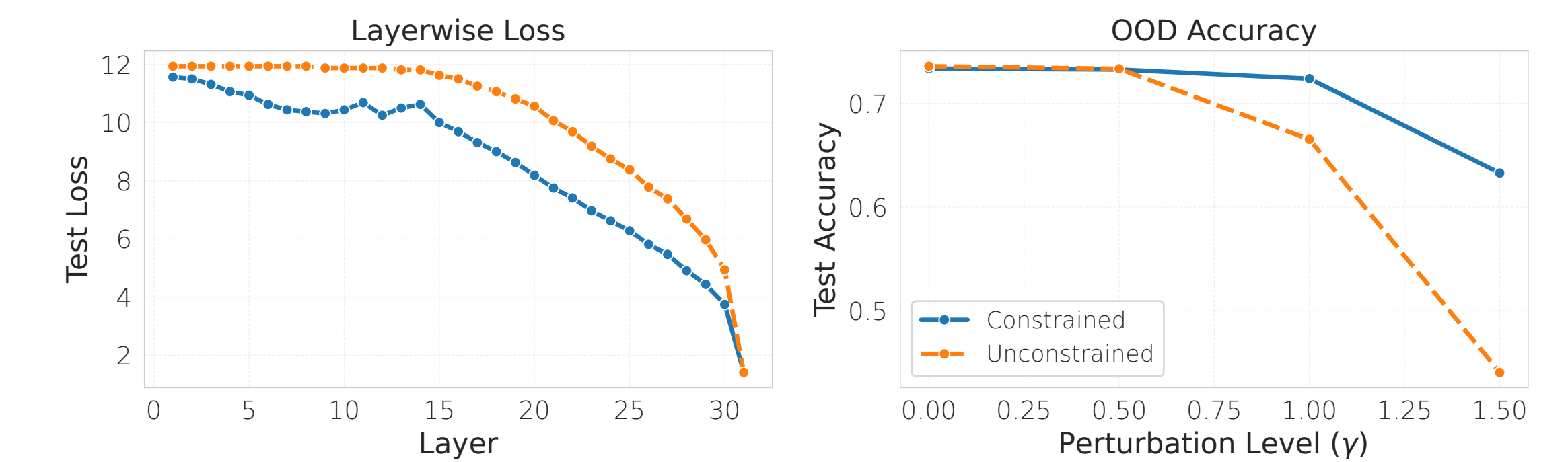


Figure 4. Constraints improve token accuracy (Left: $\downarrow$, Right: $\uparrow$).

- AlpacaEval (constrained vs unconstrained):** $\gamma = 0$: 50% win rate $\rightarrow$ preserves ID. $\gamma = 1.5$: 69.9% win rate $\rightarrow$ robust OOD.

Unrolled video denoisers generalize to OOD noise

- Task:** Reconstruct video $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_T]$ from noisy $\tilde{\mathbf{X}}$,

$$f(\mathbf{X}, \Phi(\tilde{\mathbf{X}}; \mathbf{T})) = \frac{1}{2} \sum_{t=1}^T \|\mathbf{x}_t - [\Phi(\tilde{\mathbf{X}}; \mathbf{T})]_t\|_2^2$$

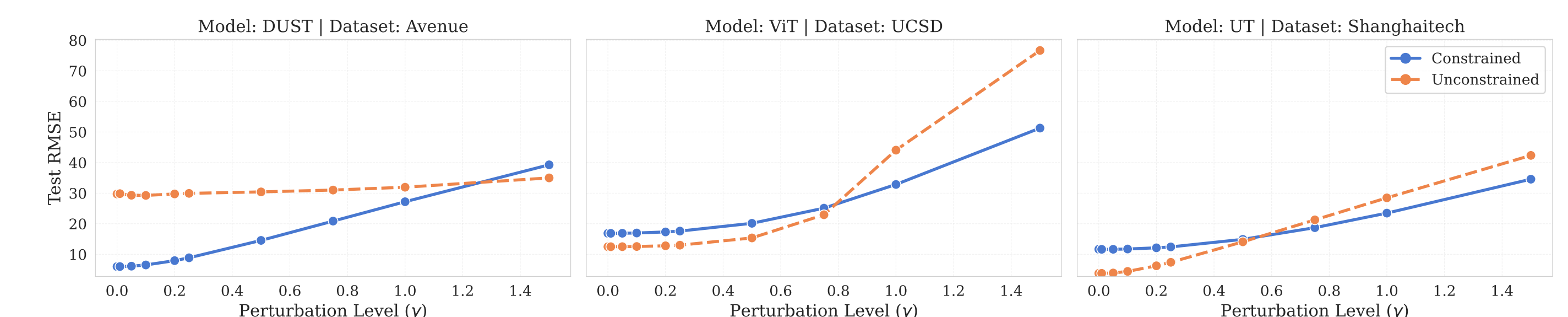


Figure 5. Constraints improve robustness. RMSE vs. test perturbation γ ($\downarrow$ lower is better).

Ablations: Step size, dual LR, training perturbations

- Low sensitivity to dual LR η_2 ; OOD accuracy increases with step size α .

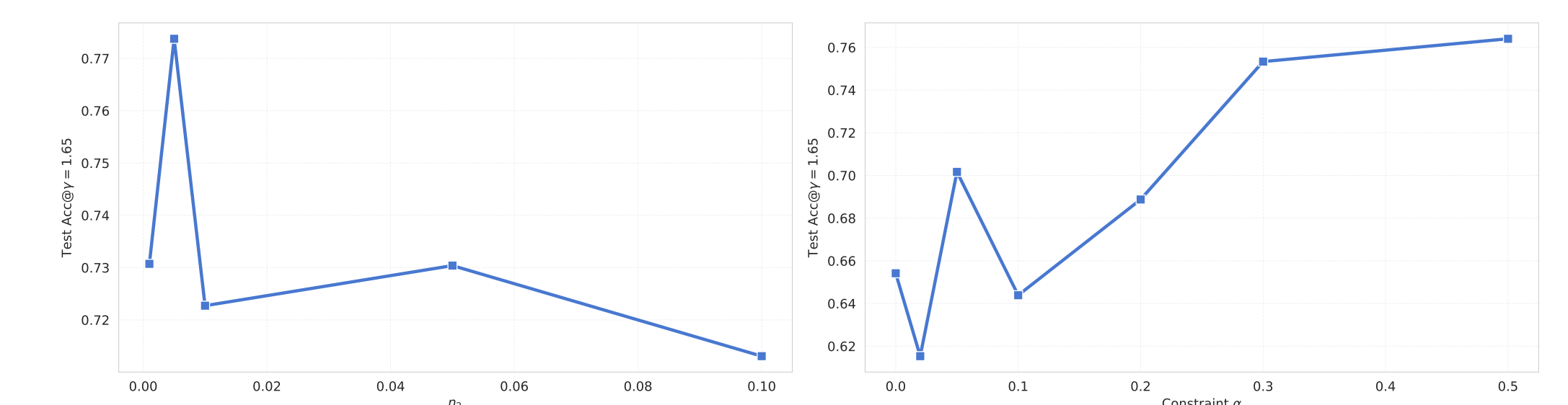


Figure 6. Dual learning rate η_2 (left) and step size α (right) ablations.

- Constraints compose with standard robustness techniques (noise injection).

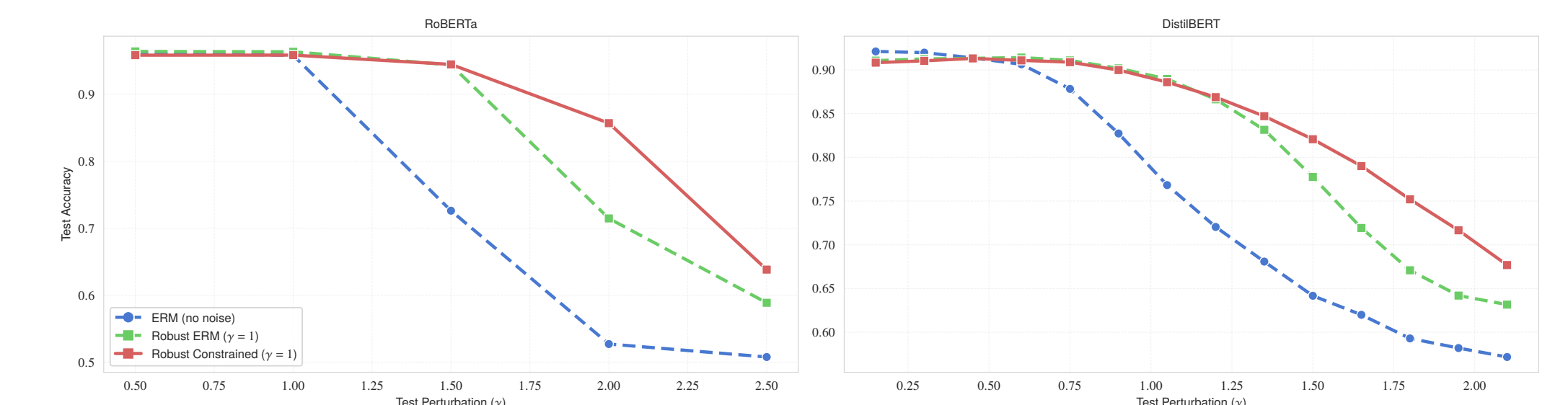


Figure 7. Comparison of baseline, robust ERM ($\gamma = 1$), and robust + constraints.