

Learning Single-Index Models In-Context

In-Context Learning (ICL): Transformers can construct various regression algorithms at inference time.

Prompt : $\underbrace{\mathbf{x}_1 \rightarrow \mathbf{y}_1, \dots, \mathbf{x}_N \rightarrow \mathbf{y}_N}_{\text{in-context examples}}, \underbrace{\mathbf{x} \rightarrow ?}_{\text{query}},$ where $y_i = f_*(\mathbf{x}_i) + (\text{noise})$

- Models need to estimate f_* **without any parameter updates**.
 - Limited **statistical** and **optimization** guarantees, especially for **nonlinear targets**

Can transformers leverage **statistically efficient** mechanisms like feature learning, even at inference time?

Problem Setting: Single-Index Target Functions

$$f_*(\mathbf{x}_i) = \sigma_*(\langle \mathbf{x}_i, \boldsymbol{\beta} \rangle), \quad \boldsymbol{\beta} \in \mathbb{R}^d; \quad \mathbf{x}_i \sim \mathcal{N}(0, \mathbf{I}_d).$$

$$\sigma_*(z) = \sum_{j=\text{ie}(\sigma_*)}^{\text{deg}(\sigma_*)} \alpha_j^* \text{He}_j(z) \quad (\text{He}_j : j\text{-th Hermite polynomial}).$$

$\rightarrow$ **degree** $\text{deg}(\sigma_*)$ and **information exponent** $\text{ie}(\sigma_*)$ ($\text{ie} \leq \text{deg}$)

- f_* changes only along the feature vector $\boldsymbol{\beta}$: A **statistical testbed for feature learning**

Sample complexity lower bound for regression on test prompt

Kernel methods: $N \gtrsim d^{\Theta(\text{deg}(\sigma_*))}$ (feature learning not available) [GMMM21,DWY21]

CSQ algorithms: $N \gtrsim d^{\Theta(\text{ie}(\sigma_*))}$ (e.g. one-step GD on NN [DLS22])

SQ algorithms: $N \gtrsim d^{\Theta(\text{ge}(\sigma_*))}$ (e.g. batch-reuse SGD on NN [ADK+24,LOSW24])

- (S)GD on NNs enables feature learning: Advantage over non-adaptive kernel methods
- Generative exponent** $\text{ge}(\sigma_*) := \min_{h \in L^2} \text{ie}(h \circ \sigma_*) \leq \text{ie}(\sigma_*)$: Nonlinear label transformations introduce non-CSQ terms, enabling further efficiency.

- Prior work on ICL:** Pretrained transformers can implement **kernel methods** over test prompts, requiring $N \gtrsim r^{\Theta(\text{deg}(\sigma_*))}$ in-context examples (assuming $\boldsymbol{\beta}$ is drawn from a rank- r linear subspace) [OSSW24].
 - Suboptimal inference-time sample complexity

Question 1: Can pretrained transformers **perform feature learning at inference time**, adapting to task-specific $\boldsymbol{\beta}$?

Question 2: How **statistically efficient** can transformers be at inference time? Can they surpass the barriers of **kernel** and **CSQ**?

Main Result: Feature Learning In-Context

Model: One-layer transformer (**softmax attention** + **nonlinear MLP**)

$$f_{\text{TF}}(\mathbf{X}_{1:N}, \mathbf{y}_{1:N}, \mathbf{x}; \mathbf{W}^{KQ}, \mathbf{W}^{FV}, \mathbf{a}, \mathbf{b}) \\ = \mathbf{a}^\top \text{ReLU}(\mathbf{W}^{FV} \mathbf{E} \cdot \text{Softmax}(\text{Mask}(\rho^{-1} \cdot \mathbf{E}^\top \mathbf{W}^{KQ} \mathbf{E})) + \mathbf{b})$$

where $\mathbf{E} = \begin{bmatrix} \mathbf{x}_1 \cdots \mathbf{x}_N \mathbf{x} \\ y_1 \cdots y_N 1 \end{bmatrix}$

Pretraining Algorithm: (under parametrization $\mathbf{W}^{KQ} = \begin{bmatrix} \Gamma & 0 \\ 0^\top & 1 \end{bmatrix}$, $\mathbf{W}^{FV} = [O \ v]$)

- One gradient descent step on the attention matrix $\mathbf{W}^{KQ}$.

$$\Gamma \leftarrow \Gamma - \eta_1 \nabla_\Gamma \left[(2T_1)^{-1} \sum_{t=1}^{T_1} (f_{\text{TF}}(\mathbf{X}_{1:N_1}^t, \mathbf{y}_{1:N_1}^t, \mathbf{x}^t) - y^t)^2 + \frac{\lambda_1}{2} \|\Gamma\|_F^2 \right]$$
- Ridge regression on the MLP parameters.

$$\mathbf{a} \leftarrow \arg \min_{\mathbf{a}} (2T_2)^{-1} \sum_{t=T_1}^{T_1+T_2} (f_{\text{TF}}(\mathbf{X}_{1:N_2}^t, \mathbf{y}_{1:N_2}^t, \mathbf{x}^t) - y^t)^2 + \frac{\lambda_2}{2} \|\mathbf{a}\|^2$$

Proposition (inference-time feature learning)

For a transformer (pretrained on $T \gtrsim d^{\Theta(\text{ie}(\sigma_*))}$ tasks), the softmax attention output is approximated by

$$C_a + C_b \left(\frac{\langle \mathbf{x}, \boldsymbol{\beta} \rangle}{\sqrt{r}} \right)^{\text{ge}(\sigma_*)}, \quad \text{where } \boldsymbol{\beta} \text{ is randomly drawn from a rank-} r \leq d \text{ linear subspace.}$$

*the approximation holds **without any parameter updates for each test prompt**

Result 1: Softmax attention modules can identify the task-specific feature vector $\boldsymbol{\beta}$ (**inference-time feature learning**)

Theorem (sample complexity bound)

The pretrained transformer achieves $o_d(1)$ estimation error at inference time for f_* , provided that the number of in-context examples N satisfies $N_{\text{test}} \gtrsim r^{\Theta(\text{ge}(\sigma_*))}$.

Result 2: Nonlinear transformers can **surpass the statistical barriers** of both **kernel** methods and **CSQ** algorithms

Sample Complexity Rates

[OSSW24]

Pretraining: $d^{\Theta(\text{ie}(\sigma_*))}$

Inference: $r^{\Theta(\text{deg}(\sigma_*))}$

This work

Pretraining: $d^{\Theta(\text{ie}(\sigma_*))}$

Inference: $r^{\Theta(\text{ge}(\sigma_*))}$

- End-to-end **statistical and optimization guarantees**
- Sample complexity depends only on the inner dimension $r \rightarrow$ **adaptive to low-dimensional prior** [OSSW24].

Technical Overview

Fact: Correlational signal $\mathbb{E}[y\phi(\mathbf{x})]$ in CSQ gets weaker as $\text{ie}(\sigma_*)$ increases...

Obs: Nonlinear label transformations in softmax **reduce the information exponent**.

- Model output is $f_{\text{TF}} = \sum_{j=1}^m a_j \sigma(v_j g(\mathbf{X}, \mathbf{y}, \mathbf{x}; \Gamma) + b_j)$ where the attention output g is

$$g(\mathbf{X}, \mathbf{y}, \mathbf{x}; \Gamma) := \frac{N^{-1} \sum_{i=1}^N y_i \exp(y_i / \rho) \exp(\mathbf{x}_i^\top \Gamma \mathbf{x} / \rho)}{N^{-1} \sum_{i=1}^N \underbrace{\exp(y_i / \rho) \exp(\mathbf{x}_i^\top \Gamma \mathbf{x} / \rho)}_{\text{Correlation between exponentially-transformed labels}}}$$

Proposition

The IE of (the clipped version of) $\sigma_* \exp(\sigma_*/\rho)$ and $\exp(\sigma_*/\rho)$ is equal to $\text{ge}(\sigma_*)$.

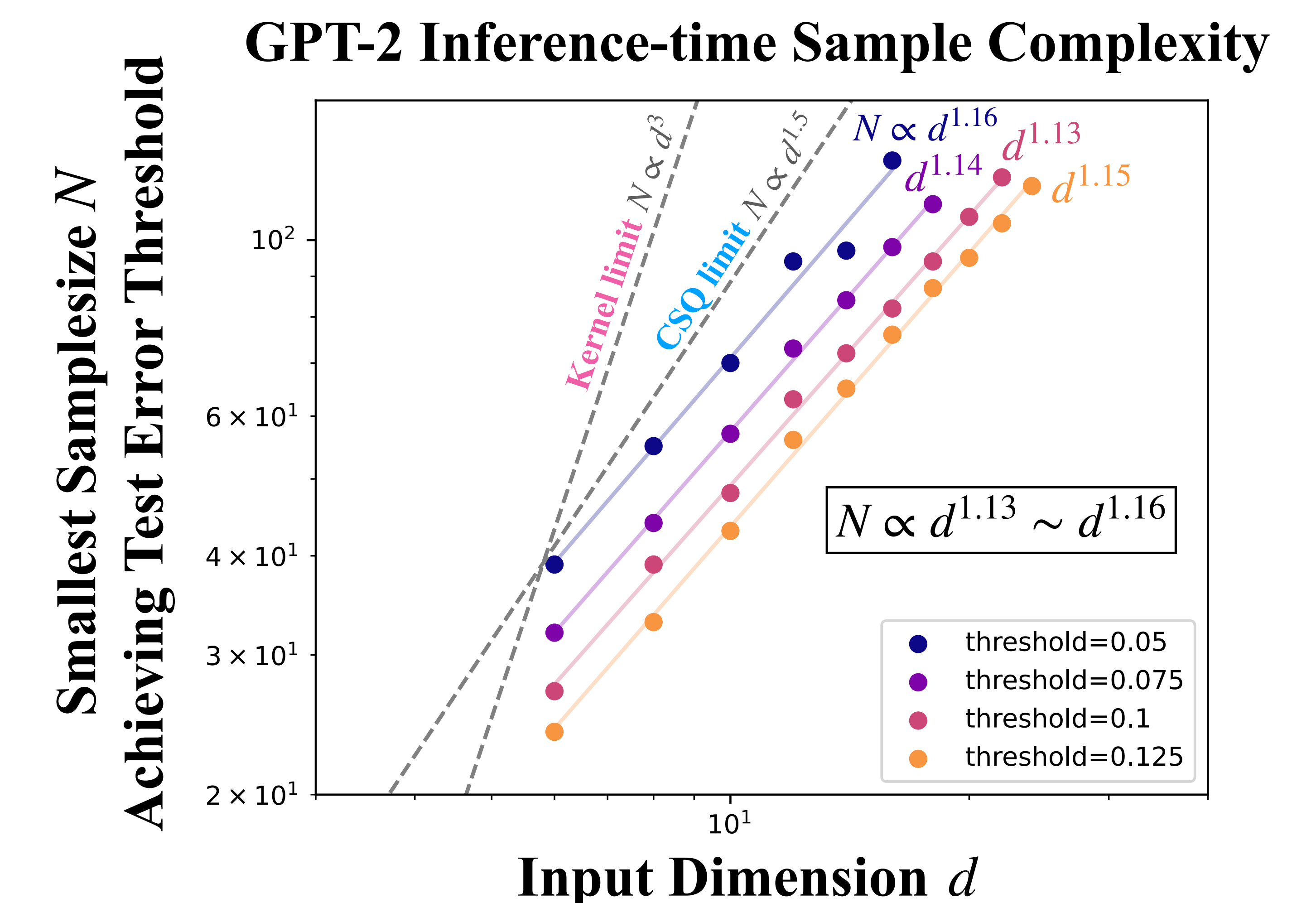
Technical Challenges to derive $\text{ge}(\sigma_*)$ -dependent sample complexity

- Softmax correlations (fractions):** Challenges in handling correlations between numerator and denominator
- Optimization guarantee:** (One GD step on Γ) $\simeq \mathbb{E}[\boldsymbol{\beta} \boldsymbol{\beta}^\top]$ (rank- r) $\rightarrow$ inference-time sample complexity $r^{\Theta(\text{ge}(\sigma_*))}$ free from d [OSSW24].

Experiment: Inference-Time Feature Learning of GPT-2

Measuring the **inference-time sample complexity** for $f_*(\mathbf{x}) = \text{He}_3(\langle \boldsymbol{\beta}, \mathbf{x} \rangle)$

Model: 6-layer GPT-2 with 4 attention heads, trained using the Adam optimizer



Result: The GPT-2 model achieves inference-time sample complexity scaling, surpassing both the kernel lower bound $\tilde{O}(d^3)$ and the CSQ lower bound $\tilde{O}(d^{1.5})$.