
Max-Margin Token Selection in Attention Mechanism

Davoud Ataee Tarzanagh
University of Pennsylvania
tarzanaq@upenn.edu

Yingcong Li Xuechen Zhang
University of California, Riverside
{yli692, xzhan394}@ucr.edu

Samet Oymak
University of Michigan
UC Riverside
oymak@umich.edu

Abstract

Attention mechanism is a central component of the transformer architecture which led to the phenomenal success of large language models. However, the theoretical principles underlying the attention mechanism are poorly understood, especially its nonconvex optimization dynamics. In this work, we explore the seminal softmax-attention model $f(X) = \langle X\mathbf{v}, \text{softmax}(XW\mathbf{p}) \rangle$, where X is the token sequence and $(\mathbf{v}, W, \mathbf{p})$ are trainable parameters. We prove that running gradient descent on $\mathbf{p}$, or equivalently W , converges in direction to a max-margin solution that separates *locally-optimal* tokens from non-optimal ones. This clearly formalizes attention as an optimal token selection mechanism. Remarkably, our results are applicable to general data and precisely characterize *optimality* of tokens in terms of the value embeddings $X\mathbf{v}$ and problem geometry. We also provide a broader regularization path analysis that establishes the margin maximizing nature of attention even for nonlinear prediction heads. When optimizing $\mathbf{v}$ and $\mathbf{p}$ simultaneously with logistic loss, we identify conditions under which the regularization paths directionally converge to their respective hard-margin SVM solutions where $\mathbf{v}$ separates the input features based on their labels. Interestingly, the SVM formulation of $\mathbf{p}$ is influenced by the support vector geometry of $\mathbf{v}$. Finally, we verify our theoretical findings via numerical experiments and provide insights.

1 Introduction

Since its introduction in the seminal work [1], attention mechanism has played an influential role in advancing natural language processing, and more recently, large language models [2, 3, 4, 5]. Initially introduced for encoder-decoder RNN architectures, attention allows the decoder to focus on the most relevant parts of the input sequence, instead of relying solely on a fixed-length hidden state. Attention mechanism has taken the center stage in the transformers [6], where the self-attention layer – which calculates softmax similarities between input tokens – serves as the backbone of the architecture. Since their inception, transformers have revolutionized natural language processing, from models like BERT [7] to ChatGPT [8], and have also become the architecture of choice for foundation models [9] addressing diverse challenges in generative modeling [3, 10], computer vision [11, 12], and reinforcement learning [13, 14, 15].

The prominence of the attention mechanism motivates a fundamental theoretical understanding of its role in optimization and learning. While it is well-known that attention enables the model to focus on the relevant parts of the input sequence, the precise mechanism by which this is achieved is far from clear. To this end, we ask

Q: What are the optimization dynamics and inductive biases of the attention mechanism?

We study this question using the fundamental attention model $f(X) = \langle X\mathbf{v}, \mathbb{S}(XW^\top \mathbf{p}) \rangle$. Here, X is the sequence of input tokens, $\mathbf{v}$ is the prediction head, W is the trainable key-query weights, and $\mathbb{S}$

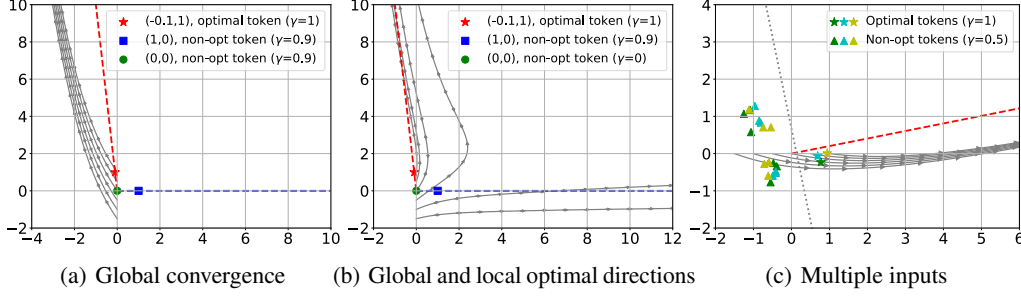


Figure 1: The convergence behavior of the gradient descent on the attention weights $\mathbf{p}$ using the logistic loss in (ERM). The arrows ($\longrightarrow$) represent trajectories from different initializations. Here, $(- - -)$ and $(- \cdot - \cdot)$ denote the globally- and locally-optimal max-margin directions (GMM, LMM). γ denotes the score of a token per Definition 1. Discussion is provided under Theorems 2 and 3.

denotes the softmax nonlinearity. For transformers, $\mathbf{p}$ corresponds to the [CLS] token or tunable prompt [16, 17, 18], whereas for RNN architectures [1], $\mathbf{p}$ corresponds to the hidden state. Given training data $(Y_i, X_i)_{i=1}^n$ with labels $Y_i \in \{-1, 1\}$ and inputs $X_i \in \mathbb{R}^{T \times d}$, we consider the empirical risk minimization with a decreasing loss function $\ell(\cdot) : \mathbb{R} \rightarrow \mathbb{R}$,

$$\mathcal{L}(\mathbf{v}, \mathbf{p}, \mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot f(X_i)), \text{ where } f(X_i) = \mathbf{v}^\top X_i^\top \mathbb{S}(X_i \mathbf{W}^\top \mathbf{p}). \quad (1)$$

At a high-level, this work establishes fundamental equivalences between the optimization trajectories of (1) and hard-margin SVM problems. Our main contributions are as follows:

- **Optimization geometry of attention (Sec 2):** We first show that gradient iterations of $\mathbf{p}$ and $\mathbf{W}$ admit a one-to-one mapping, thus we focus on optimizing $\mathbf{p}$ without losing generality. In Theorem 3, we prove that, under proper initialization:

Gradient descent on $\mathbf{p}$ converges in direction to a max-margin solution – namely (ATT-SVM) – that separates locally-optimal tokens from non-optimal ones.

We call these Locally-optimal Max-Margin (LMM) directions and show that these thoroughly characterize the viable convergence directions of attention when the norm of its weights grows to infinity. We also identify conditions under which (algorithm-independent) regularization path and gradient descent path converge to Globally-optimal Max-Margin (GMM) direction in Theorems 1 and 2, respectively. A central feature of our results is precisely quantifying *optimality* in terms of *token scores* $\gamma_i = Y \cdot \mathbf{v}^\top \mathbf{x}_i$ where $\mathbf{x}_i$ is the i^{th} token of the input sequence X . *Locally-optimal* tokens are those with higher scores than their nearest neighbors determined by the SVM solution. These are illustrated in Figure 1.

- **Optimize attention $\mathbf{p}$ and prediction-head $\mathbf{v}$ jointly (Sec 3):** We study the joint problem under logistic loss function. We use regularization path analysis where (ERM) is solved under ridge constraints and we study the solution trajectory as the constraints are relaxed. Since the problem is linear in $\mathbf{v}$, if the attention features $\mathbf{x}_i^{\text{att}} = X_i^\top \mathbb{S}(X_i \mathbf{W}^\top \mathbf{p})$ are separable based on their labels Y_i , $\mathbf{v}$ would implement a max-margin classifier. Building on this, we prove that $\mathbf{p}$ and $\mathbf{v}$ converges to their respective max-margin solutions under proper geometric conditions (Theorem 5). Relaxing these conditions, we obtain a more general solution where margin constraints on $\mathbf{p}$ are relaxed on the inputs whose attention features are not support vectors of $\mathbf{v}$ (Theorem 6). Figure 3 illustrates these outcomes.

The next section introduces the preliminary concepts, Section 4 presents numerical experiments¹, Section 5 discusses related literature, and Section 6 highlights limitations and future work.

1.1 Preliminaries

Notations. For any integer $N \geq 1$, let $[N] := \{1, \dots, N\}$. We use lower-case and upper-case bold letters (e.g. $\mathbf{a}$ and $\mathbf{A}$) to represent vectors and matrices, respectively. The entries of $\mathbf{a}$ are denoted as a_i . We use $\bar{\sigma}(\mathbf{A})$ to denote the maximum singular value of $\mathbf{A}$. We denote the minimum of two numbers a, b as $a \wedge b$, and the maximum as $a \vee b$. Big-O notation $\mathcal{O}(\cdot)$ hides the universal constants. Throughout, we will use $\mathcal{L}(\mathbf{p})$ and $\mathcal{L}(\mathbf{v}, \mathbf{p})$ to denote Objective (1) with fixed $(\mathbf{v}, \mathbf{W})$ and $\mathbf{W}$, respectively.

Optimization. Given an objective function $\mathcal{L} : \mathbb{R}^d \rightarrow \mathbb{R}$ and an ℓ_2 -norm bound R , define the regularized solution as

¹The code for experiments can be found at https://github.com/ucr-optml/max_margin_attention.

$$\bar{\mathbf{p}}(R) := \arg \min_{\|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p}). \quad (2)$$

Regularization path – the evolution of $\bar{\mathbf{p}}(R)$ as R grows – is known to capture the spirit of gradient descent as the ridge constraint R provides a proxy for the number of gradient descent iterations. For instance, [19, 20, 21] study the implicit bias of logistic regression and rigorously connect the directional convergence of regularization path (i.e. $\lim_{R \rightarrow \infty} \bar{\mathbf{p}}(R)/R$) and gradient descent. For gradient descent, we assume the objective $\mathcal{L}(\mathbf{p})$ is smooth and describe the gradient descent process as

$$\mathbf{p}(t+1) = \mathbf{p}(t) - \eta(t) \nabla \mathcal{L}(\mathbf{p}(t)), \quad (3)$$

where $\eta(t)$ is the stepsize at time t and $\nabla \mathcal{L}(\mathbf{p}(t))$ is the gradient of $\mathcal{L}$ at $\mathbf{p}(t)$.

Attention in Transformers. Next, we will discuss the connection between our model and the attention mechanism used in transformers. Our exposition borrows from [17], where the authors analyze the same attention model using gradient-based techniques on specific contextual datasets.

• **Self-attention** is the core building block of transformers [6]. Given an input consisting of T tokens $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_T]^\top \in \mathbb{R}^{T \times d}$, self-attention with key-query matrix $\mathbf{W} \in \mathbb{R}^{d \times d}$, and value matrix $\mathbf{V} \in \mathbb{R}^{d \times v}$, the self-attention model is defined as follows:

$$f_{sa}(\mathbf{X}) = \mathbb{S}(\mathbf{X}\mathbf{W}\mathbf{X}^\top)\mathbf{X}\mathbf{V}. \quad (4)$$

Here, $\mathbb{S}(\cdot)$ is the softmax nonlinearity that applies row-wise on the similarity matrix $\mathbf{X}\mathbf{W}\mathbf{X}^\top$.

• **Tunable tokens: [CLS] and prompt-tuning.** In practice, we append additional tokens to the raw input features $\mathbf{X}$: For instance, a [CLS] token is used for classification purposes [7] and prompt vectors can be appended for adapting a pretrained model to new tasks [16, 18]. Let $\mathbf{p} \in \mathbb{R}^d$ be the tunable token ([CLS] or prompt vector) and concatenate it to $\mathbf{X}$ to obtain $\mathbf{X}_p := [\mathbf{p} \ \mathbf{X}^\top]^\top \in \mathbb{R}^{(T+1) \times d}$. Consider the cross-attention features obtained from $\mathbf{X}_p$ and $\mathbf{X}$ given by

$$\begin{bmatrix} f_{cls}^\top(\mathbf{X}) \\ f_{sa}^\top(\mathbf{X}) \end{bmatrix} = \mathbb{S}(\mathbf{X}_p \mathbf{W} \mathbf{X}^\top) \mathbf{X} \mathbf{V} = \begin{bmatrix} \mathbb{S}(\mathbf{p}^\top \mathbf{W} \mathbf{X}^\top) \\ \mathbb{S}(\mathbf{X} \mathbf{W} \mathbf{X}^\top) \end{bmatrix} \mathbf{X} \mathbf{V}.$$

The beauty of cross-attention is that it isolates the contribution of $\mathbf{p}$ under the upper term $f_{cls}(\mathbf{X}) = \mathbf{V}^\top \mathbf{X}^\top \mathbb{S}(\mathbf{X} \mathbf{W} \mathbf{X}^\top) \mathbf{p} \in \mathbb{R}^v$. In this work, we use the value weights for classification, thus we set $v = 1$, and denote $\mathbf{v} = \mathbf{V} \in \mathbb{R}^d$. This brings us to our attention model of interest:

$$f(\mathbf{X}) = \mathbf{v}^\top \mathbf{X}^\top \mathbb{S}(\mathbf{K} \mathbf{p}), \quad \text{where } \mathbf{K} = \mathbf{X} \mathbf{W}^\top. \quad (5)$$

Here, $(\mathbf{v}, \mathbf{W}, \mathbf{p})$ are the tunable model parameters and $\mathbf{K}$ is the key embeddings. Note that $\mathbf{W}$ and $\mathbf{p}$ are playing the same role within softmax, thus, it is intuitive that they exhibit similar optimization dynamics. Confirming this, the next lemma shows that gradient iterations of $\mathbf{p}$ (after setting $\mathbf{W} \leftarrow \text{Identity}$) and $\mathbf{W}$ admit a one-to-one mapping.

Lemma 1 Fix $\mathbf{u} \in \mathbb{R}^d \setminus \{\mathbf{0}\}$. Let $\psi : \mathbb{R}^d \rightarrow \mathbb{R}$ and $\ell : \mathbb{R} \rightarrow \mathbb{R}$ be differentiable functions. On the same training data $(Y_i, \mathbf{X}_i)_{i=1}^n$, define $\tilde{\mathcal{L}}(\mathbf{p}) := 1/n \sum_{i=1}^n \ell(Y_i \cdot \psi(\mathbf{X}_i^\top \mathbb{S}(\mathbf{X}_i \mathbf{p})))$ and $\mathcal{L}(\mathbf{W}) := 1/n \sum_{i=1}^n \ell(Y_i \cdot \psi(\mathbf{X}_i^\top \mathbb{S}(\mathbf{X}_i \mathbf{W}^\top \mathbf{u})))$. Consider the gradient descent iterations on $\mathbf{p}$ and $\mathbf{W}$ with initial values $\mathbf{p}(0)$ and $\mathbf{W}(0) = \mathbf{u} \mathbf{p}(0)^\top / \|\mathbf{u}\|^2$ and stepsizes η and $\eta/\|\mathbf{u}\|^2$, respectively:

$$\mathbf{p}(t+1) = \mathbf{p}(t) - \eta \nabla \tilde{\mathcal{L}}(\mathbf{p}(t)),$$

$$\mathbf{W}(t+1) = \mathbf{W}(t) - \frac{\eta}{\|\mathbf{u}\|^2} \nabla \mathcal{L}(\mathbf{W}(t)).$$

We have that $\mathbf{W}(t) = \mathbf{u} \mathbf{p}(t)^\top / \|\mathbf{u}\|^2$ for all $t \geq 0$.

This lemma directly characterizes the optimization dynamics of $\mathbf{W}$ through the dynamics of $\mathbf{p}$, allowing us to reconstruct $\mathbf{W}$ from $\mathbf{p}$ using their gradient iterations. Therefore, we will fix $\mathbf{W}$ and concentrate on optimizing $\mathbf{p}$ in Section 2 and the joint optimization of $(\mathbf{v}, \mathbf{p})$ in Section 3.

Problem definition: Throughout, $(Y_i, \mathbf{X}_i)_{i=1}^n$ denotes training dataset where $Y_i \in \{-1, 1\}$ and $\mathbf{X}_i \in \mathbb{R}^{T \times d}$. We denote the key embeddings of $\mathbf{X}_i$ via $\mathbf{K}_i = \mathbf{X}_i \mathbf{W}^\top$ and explore the training risk

$$\mathcal{L}(\mathbf{v}, \mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot \mathbf{v}^\top \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p})). \quad (\text{ERM})$$

Importantly, our results apply to general tuples $(Y_i, \mathbf{X}_i, \mathbf{K}_i)$ and do not assume that $(\mathbf{X}_i, \mathbf{K}_i)$ are tied via $\mathbf{W}$. Finally, the t^{th} tokens of $\mathbf{X}_i, \mathbf{K}_i$ are denoted by $\mathbf{x}_{it}, \mathbf{k}_{it} \in \mathbb{R}^d$, respectively, for $t \in [T]$.

The highly nonlinear and nonconvex nature of the softmax operation makes the training problem in (ERM) a challenging nonconvex optimization problem for $\mathbf{p}$, even with a fixed $\mathbf{v}$. In the next section, we will introduce a set of assumptions to demonstrate the global and local convergence of gradient descent for margin maximization in the attention mechanism.

2 Global and Local Margin Maximization with Attention

In this section, we present the main results of this paper (Theorems 2 and 3) by examining the implicit bias of gradient descent on learning $\mathbf{p} \in \mathbb{R}^d$ given a fixed choice of $\mathbf{v} \in \mathbb{R}^d$. Notably, our results apply to general *decreasing loss functions without requiring convexity*. This generality is attributed to margin maximization arising from the exponentially-tailed nature of softmax within attention, rather than ℓ . We maintain the following assumption on the loss function throughout this section.

Assumption A (Well-behaved Loss) *Over any bounded interval: (1) $\ell : \mathbb{R} \rightarrow \mathbb{R}$ is strictly decreasing. (2) ℓ' is M_0 -Lipschitz continuous and $|\ell'(u)| \leq M_1$.*

Assumption A includes many common loss functions, including the logistic loss $\ell(u) = \log(1 + e^{-u})$, exponential loss $\ell(u) = e^{-u}$, and correlation loss $\ell(u) = -u$. Assumption A implies that $\mathcal{L}(\mathbf{p})$ is L_p -smooth (see Lemma 6 in Supplementary), where

$$L_p := \frac{1}{n} \sum_{i=1}^n \left(M_0 \|\mathbf{v}\|^2 \|\mathbf{W}\|^2 \|X_i\|^4 + 3M_1 \|\mathbf{v}\| \|\mathbf{W}\|^2 \|X_i\|^3 \right). \quad (6)$$

We now introduce a convex hard-margin SVM problem that separates one token of the input sequence from the rest, jointly solved over all inputs. We will show that this problem captures the optimization properties of softmax-attention. Fix indices $\alpha = (\alpha_i)_{i=1}^n$ and consider

$$\mathbf{p}^{mm}(\alpha) = \arg \min_{\mathbf{p}} \|\mathbf{p}\| \quad \text{subject to} \quad \min_{i \neq \alpha_i} \mathbf{p}^\top (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}) \geq 1, \quad \text{for all } 1 \leq i \leq n. \quad (\text{ATT-SVM})$$

Note that existence of $\mathbf{p}^{mm}(\alpha)$ implies the separability of tokens α from the others. Specifically, choosing direction $\mathbf{p}^{mm}(\alpha)$ will exactly select tokens $(x_{i\alpha_i})_{i=1}^n$ at the attention output for each input sequence, that is, $\lim_{R \rightarrow \infty} X_i^\top \mathbb{S}(R \cdot \mathbf{K}_i \mathbf{p}^{mm}(\alpha)) = x_{i\alpha_i}$. We are now ready to introduce our main results that characterize the global and local convergence of the attention weights $\mathbf{p}$ via (ATT-SVM).

2.1 Global convergence of the attention weights $\mathbf{p}$

We first identify the conditions that guarantee the global convergence of gradient descent for $\mathbf{p}$. The intuition is that, in order for attention to exhibit implicit bias, the softmax nonlinearity should be forced to select the *optimal token within each input sequence*. Fortunately, the optimal tokens that achieve the smallest training objective under decreasing loss function $\ell(\cdot)$ have a clear definition.

Definition 1 (Token Scores, Optimality & GMM) *The score of token x_{it} of input X_i is defined as $\gamma_{it} := Y_i \cdot \mathbf{v}^\top x_{it}$. The optimal tokens for input X_i are those tokens with highest scores given by*

$$\text{opt}_i \in \arg \max_{i \in [T]} \gamma_{it}.$$

Globally-optimal max-margin (GMM) direction is defined as the solution of (ATT-SVM) with optimal indices $(\text{opt}_i)_{i=1}^n$ by $\mathbf{p}^{mm}$.*

It is worth noting that score definition simply uses the *value embeddings* $\mathbf{v}^\top x_{it}$ of the tokens. Note that multiple tokens within an input might attain the same score, thus opt_i or $\mathbf{p}^{mm*}$ may not be unique. The theorem below provides our regularization path guarantee on the global convergence of attention.

Theorem 1 (Regularization Path) *Suppose Assumption A on the loss function holds, and for all $i \in [n]$ and $t \neq \text{opt}_i$, the scores obey $\gamma_{it} < \gamma_{i\text{opt}_i}$. Then, the regularization path $\tilde{\mathbf{p}}(R) = \arg \min_{\|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p})$ converges to the GMM direction i.e. $\lim_{R \rightarrow \infty} \tilde{\mathbf{p}}(R)/R = \mathbf{p}^{mm*}/\|\mathbf{p}^{mm*}\|$.*

Theorem 1 shows that as the regularization strength R increases towards the ridgeless problem $\min_{\mathbf{p}} \mathcal{L}(\mathbf{p})$, the optimal direction $\tilde{\mathbf{p}}(R)$ aligns more closely with the max-margin solution $\mathbf{p}^{mm*}$. Since this theorem allows for arbitrary token scores, it demonstrates that *max-margin token separation is an essential feature of the attention mechanism*. In fact, it is a corollary of Theorem 8, which applies to

the generalized model $f(X) = \psi(X^\top \mathbb{S}(XW^\top \mathbf{p}))$ and accommodates multiple optimal tokens per input. However, while regularization path analysis captures the global behavior, gradient descent lacks general global convergence guarantees. In Section 2.2, we show that due to the nonconvex landscape and softmax nonlinearity, gradient descent often converges to local optima. We first establish that when (ERM) is trained with gradient descent, the norm of the parameters will diverge. For the restrictive setting of $n = 1$, gradient descent also exhibits a global convergence guarantee.

Assumption B For all $i \in [n]$ and $t, \tau \neq \text{opt}_i$, the scores per Definition 1 obey $\gamma_{it} = \gamma_{i\tau} < \gamma_{i\text{opt}_i}$.

Theorem 2 (Global Convergence of Gradient Descent) Suppose Assumption A on the loss function ℓ and Assumption B on the tokens' score hold. Then, the gradient descent iterates $\mathbf{p}(t+1) = \mathbf{p}(t) - \eta \nabla \mathcal{L}(\mathbf{p}(t))$ on (ERM), with the stepsize $\eta \leq 1/L_p$ and any starting point $\mathbf{p}(0)$ satisfy $\lim_{t \rightarrow \infty} \|\mathbf{p}(t)\| = \infty$. If $n = 1$, we also have $\lim_{t \rightarrow \infty} \mathbf{p}(t)/\|\mathbf{p}(t)\| = \mathbf{p}^{\text{mm}\star}/\|\mathbf{p}^{\text{mm}\star}\|$.

Theorem 2 shows that gradient descent will diverge in norm, and when $n = 1$, the normalized predictor $\mathbf{p}(t)/\|\mathbf{p}(t)\|$ converges towards $\mathbf{p}^{\text{mm}\star}$, the separator of the globally optimal token. While $n = 1$ is a stringent condition, this requirement is in fact tight as discussed in Appendix E. To illustrate this theorem, we have conducted synthetic experiments. Let us first explain the setup used in Figure 1. We set $d = 3$ as the dimension, with each token having three entries $\mathbf{x} = [x_1, x_2, x_3]$. We reserve the first two coordinates as key embeddings $\mathbf{k} = [x_1, x_2, 0]$ by setting $\mathbf{W} = \text{diag}([1, 1, 0])$. This is what we display in our figures as token positions. Finally, in order to assign scores to the tokens we use the last coordinate by setting $\mathbf{v} = [0, 0, 1]$. This way score becomes $Y \cdot \mathbf{v}^\top \mathbf{x} = Y \cdot x_3$, allowing us to assign any score (regardless of key embedding).

In Figure 1(a), the gray paths represent gradient descent trajectories from different initializations. The points (0, 0) and (1, 0) correspond to non-optimal tokens, while the point (-0.1, 1) represents the optimal token. Notably, gradient descent iterates with various starting points converge towards the direction of the max-margin solution $\mathbf{p}^{\text{mm}\star}$ (depicted by - - -). Moreover, as the iteration count t increases, the inner product $\langle \mathbf{p}(t)/\|\mathbf{p}(t)\|, \mathbf{p}^{\text{mm}\star}/\|\mathbf{p}^{\text{mm}\star}\| \rangle$ consistently increases. Figure 1(c) also depicts the directional convergence of gradient descent from various initializations on multiple inputs, with the gray dotted line representing the separating hyperplane. These emphasize the gradual alignment between the evolving predictor and the max-margin solution throughout the optimization.

Lemma 2 Suppose for all $i \in [n]$ and $t \neq \text{opt}_i$, $Y_i = 1$ and $\gamma_{it} < \gamma_{i\text{opt}_i}$. Also assume $\mathbf{W} \in \mathbb{R}^{d \times d}$ is full-rank. Then $\mathbf{p}^{\text{mm}\star}$ exists – i.e. (ATT-SVM) is feasible for optimal indices $\alpha_i \leftarrow \text{opt}_i$.

2.2 Local convergence of the attention weights $\mathbf{p}$

Theorem 2 on the global convergence of gradient descent serves as a prelude to the general behavior of the optimization. Once we relax Assumption B by allowing for arbitrary token scores, we will show that $\mathbf{p}$ can converge (in direction) to a locally-optimal solution. However, this locally-optimal solution is still characterized in terms of (ATT-SVM) which separates *locally-optimal* tokens from the rest. Our theory builds on two new concepts: locally-optimal tokens and neighbors of these tokens.

Definition 2 (SVM-Nighbor and Locally-Optimal Tokens) Fix token indices $\alpha = (\alpha_i)_{i=1}^n$ for which (ATT-SVM) is feasible to obtain $\mathbf{p}^{\text{mm}} = \mathbf{p}^{\text{mm}}(\alpha)$. Consider tokens $\mathcal{T}_i \subset [T]$ such that $(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{p}^{\text{mm}} = 1$ for all $t \in \mathcal{T}_i$. We refer to $\mathcal{T}_i$ as SVM-neighbors of $\mathbf{k}_{i\alpha_i}$. Additionally, tokens with indices $\alpha = (\alpha_i)_{i=1}^n$ are called locally-optimal if for all $i \in [n]$ and $t \in \mathcal{T}_i$ scores per Definition 1 obey $\gamma_{i\alpha_i} > \gamma_{it}$. Associated $\mathbf{p}^{\text{mm}}$ is called a locally-optimal max-margin (LMM) direction.

To provide a basis for discussing local convergence, we provide some preliminary definitions regarding cones. For a given $\mathbf{q}$ and a scalar $\mu > 0$, we define $\text{cone}_\mu(\mathbf{q})$ as the set of vectors $\mathbf{p} \in \mathbb{R}^d$ such that the correlation coefficient between $\mathbf{p}$ and $\mathbf{q}$ is at least $1 - \mu$:

$$\text{cone}_\mu(\mathbf{q}) := \left\{ \mathbf{p} \in \mathbb{R}^d \mid \left\langle \frac{\mathbf{p}}{\|\mathbf{p}\|}, \frac{\mathbf{q}}{\|\mathbf{q}\|} \right\rangle \geq 1 - \mu \right\}. \quad (7)$$

Given $R > 0$, the intersection of $\text{cone}_\mu(\mathbf{q})$ and the set $\{\mathbf{p} \in \mathbb{R}^d \mid \|\mathbf{p}\| \geq R\}$ is denoted as $C_{\mu,R}(\mathbf{q})$:

$$C_{\mu,R}(\mathbf{q}) := \text{cone}_\mu(\mathbf{q}) \cap \{\mathbf{p} \in \mathbb{R}^d \mid \|\mathbf{p}\| \geq R\}. \quad (8)$$

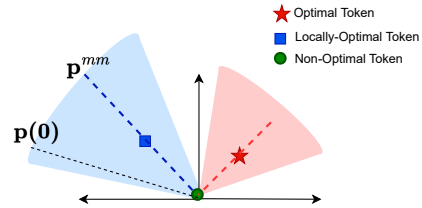


Figure 2: Gradient descent initialization $\mathbf{p}(0)$ inside the cone containing the locally-optimal solution $\mathbf{p}^{\text{mm}}$.

Next, we demonstrate the existence of parameters $\mu = \mu(\alpha) > 0$ and $R > 0$ such that when R is sufficiently large, there are no stationary points within $C_{\mu,R}(\mathbf{p}^{mm})$. Further, the gradient descent initialized within $C_{\mu,R}(\mathbf{p}^{mm})$ converges in direction to $\mathbf{p}^{mm}/\|\mathbf{p}^{mm}\|$; refer to Figure 2 for a visualization.

Theorem 3 (Local Convergence of Gradient Descent) *Suppose Assumption A on the loss function ℓ holds and assume $\alpha = (\alpha_i)_{i=1}^n$ are indices of locally-optimal tokens per Definition 2. Then, there is a constant $\mu = \mu(\alpha) \in (0, 1)$ and $R > 0$ such that $C_{\mu,R}(\mathbf{p}^{mm})$ does not contain any stationary points. Further, for any starting point $\mathbf{p}(0) \in C_{\mu,R}(\mathbf{p}^{mm})$, gradient descent iterates $\mathbf{p}(t+1) = \mathbf{p}(t) - \eta \nabla \mathcal{L}(\mathbf{p}(t))$ on (ERM) with stepsize $\eta \leq 1/L_p$ satisfies $\lim_{t \rightarrow \infty} \|\mathbf{p}(t)\| = \infty$ and $\lim_{t \rightarrow \infty} \mathbf{p}(t)/\|\mathbf{p}(t)\| = \mathbf{p}^{mm}/\|\mathbf{p}^{mm}\|$.*

To further illustrate Theorem 3, we can consider Figure 1(b) where $n = 1$ and $T = 3$. In this figure, the point (0, 0) represents the non-optimal tokens, while (1, 0) represents the locally optimal token. Additionally, the gray paths represent the trajectories of gradient descent initiated from different points. By observing the figure, we can see that gradient descent, when properly initialized, converges towards the direction of $\mathbf{p}^{mm}$ (depicted by - - -). This direction of convergence effectively separates the locally optimal tokens (1, 0) from the non-optimal token (0, 0).

2.3 Regularization paths can only converge to locally-optimal max-margin directions

An important question arises regarding whether our definition of LMM (Definition 2) encompasses all possible convergence paths of the attention mechanism when $\|\mathbf{p}\| \rightarrow \infty$. To address this, we introduce the set of LMM directions as follows:

$$\mathcal{P}^{mm} := \left\{ \frac{\mathbf{p}^{mm}(\alpha)}{\|\mathbf{p}^{mm}(\alpha)\|} \mid \alpha \text{ is locally-optimal per Definition 2} \right\}.$$

The following theorem establishes the tightness of these directions: It demonstrates that for any candidate $\mathbf{q} \notin \mathcal{P}^{mm}$, its local regularization path within an arbitrarily small neighborhood will provably not converge in the direction of $\mathbf{q}$.

Theorem 4 *Fix $\mathbf{q} \notin \mathcal{P}^{mm}$ with unit ℓ_2 norm. Assume that token scores are distinct (namely $\gamma_{it} \neq \gamma_{i\tau}$ for $t \neq \tau$) and key embeddings $\mathbf{k}_{it}$ are in general position (see Theorem 7). Fix arbitrary $\epsilon > 0, R_0 > 0$. Define the local regularization path of $\mathbf{q}$ as its (ϵ, R_0) -conic neighborhood:*

$$\bar{\mathbf{p}}(R) = \arg \min_{\mathbf{p} \in C_{\epsilon,R_0}(\mathbf{q}), \|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p}), \text{ where } C_{\epsilon,R_0}(\mathbf{q}) = \text{cone}_{\epsilon}(\mathbf{q}) \cap \{\mathbf{p} \in \mathbb{R}^d \mid \|\mathbf{p}\| \geq R_0\}. \quad (9)$$

Then, either $\lim_{R \rightarrow \infty} \|\bar{\mathbf{p}}(R)\| < \infty$ or $\lim_{R \rightarrow \infty} \bar{\mathbf{p}}(R)/\|\bar{\mathbf{p}}(R)\| \neq \mathbf{q}$. In both scenarios $\lim_{R \rightarrow \infty} \bar{\mathbf{p}}(R)/R \neq \mathbf{q}$.

The result above nicely complements Theorem 3, which states that when gradient descent is initialized above a threshold ($\|\mathbf{p}(0)\| \geq R_0$) in an LMM direction, $\|\mathbf{p}(t)\|$ diverges but the direction converges to LMM. In contrast, Theorem 4 shows that regardless of how small the cone is (in terms of angle and norm lower bound $\|\mathbf{p}\| \geq R_0$), the optimal solution path will not converge along $\mathbf{q} \notin \mathcal{P}^{mm}$.

3 Joint Convergence of Head $\mathbf{v}$ and Attention Weights $\mathbf{p}$

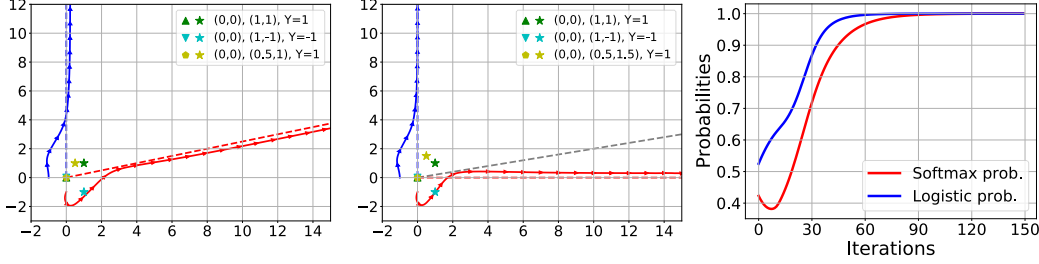
In this section, we extend the preceding results to the general case of joint optimization of head $\mathbf{v}$ and attention weights $\mathbf{p}$ using a logistic loss function. To this aim, we focus on regularization path analysis, which involves solving (ERM) under ridge constraints and examining the solution trajectory as the constraints are relaxed.

High-level intuition. Since the prediction is linear as a function of $\mathbf{v}$, logistic regression in $\mathbf{v}$ can exhibit its own implicit bias to a max-margin solution. Concretely, define the attention features $\mathbf{x}_i^p = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p})$ and define the dataset $\mathcal{D}^p = (Y_i, \mathbf{x}_i^p)_{i=1}^n$. If this dataset $\mathcal{D}^p$ is linearly separable, then fixing $\mathbf{p}$ and optimizing only $\mathbf{v}$ will converge in the direction of the standard max-margin classifier

$$\mathbf{v}^{mm} = \arg \min_{\mathbf{v} \in \mathbb{R}^d} \|\mathbf{v}\| \quad \text{subject to} \quad Y_i \cdot \mathbf{v}^\top \mathbf{r}_i \geq 1, \text{ for all } 1 \leq i \leq n, \quad (\text{SVM})$$

after setting inputs to the attention features $\mathbf{r}_i \leftarrow \mathbf{x}_i^p$ [22]. This motivates a clear question:

Under what conditions, optimizing $\mathbf{v}, \mathbf{p}$ jointly will converge to their respective max-margin solutions? We study this question in two steps. Loosely speaking: (1) We will first assume that, at the optimal



(a) All inputs are support vectors (b) (0.5, 1.5) is not a support vector (c) Probability evolutions in (a)

Figure 3: (a) and (b) Joint convergence of attention weights $\mathbf{p}$ (---) and classifier head $\mathbf{v}$ (—) to max-margin directions. (c) Averaged softmax probability evolution of optimal tokens and logistic probability evolution of output in (a).

tokens $\mathbf{x}_{i\alpha_i}, i \in [n]$ selected by $\mathbf{p}$, when solving (SVM) with $\mathbf{r}_i \leftarrow \mathbf{x}_{i\alpha_i}$, all of these tokens become support vectors of (SVM). (2) We will then relax this condition to uncover a more general implicit bias for $\mathbf{p}$ that distinguish support vs non-support vectors. Throughout, we assume that the joint problem is separable and there exists $(\mathbf{v}, \mathbf{p})$ asymptotically achieving zero training loss.

3.1 When all attention features are support vectors

In (SVM), define *label margin* to be $1/\|\mathbf{v}^{mm}\|$. Our first insight in quantifying the joint implicit bias is that, **optimal tokens** admit a natural definition: *Those that maximize the downstream label margin when selected.* This is formalized below where we assume that: (1) Selecting the token indices $\alpha = (\alpha_i)_{i=1}^n$ from each input data achieves the largest label margin. (2) The optimality of the α choice is strict in the sense that mixing other tokens will shrink the label margin in (SVM).

Assumption C (Optimal Tokens) Let $\Gamma > 0$ be the label margin when solving (SVM) with $\mathbf{r}_i \leftarrow \mathbf{x}_{i\alpha_i}$. There exists $\nu > 0$ such that for all $\mathbf{p}$, solving (SVM) with $\mathbf{r}_i \leftarrow \mathbf{x}_i^{\mathbf{p}}$ results in a label margin of at most $\Gamma - \nu \cdot \max_{i \in [n]} (1 - s_{i\alpha_i})$ where $s_i = \mathbb{S}(\mathbf{K}_i \mathbf{p})$.

Example: To gain intuition, let us fix $\mathbf{a} \in \mathbb{R}^d$ and consider the dataset obeying $\mathbf{x}_{i1} = Y_i \cdot \mathbf{a}$ and $\|\mathbf{x}_{it}\| < \|\mathbf{a}\|$ for all $t \geq 2$ and all $i \in [n]$. For this dataset, we can choose $\alpha_i = 1$, $\mathbf{v}^{mm} = \mathbf{a}/\|\mathbf{a}\|^2$, $\Gamma = 1/\|\mathbf{v}^{mm}\| = \|\mathbf{a}\|$ and $\nu = \|\mathbf{a}\| - \sup_{i \in [n], t \geq 2} \|\mathbf{x}_{it}\|$.

Theorem 5 Consider the ridge-constrained solutions $(\mathbf{v}_r, \mathbf{p}_R)$ of (ERM) defined as

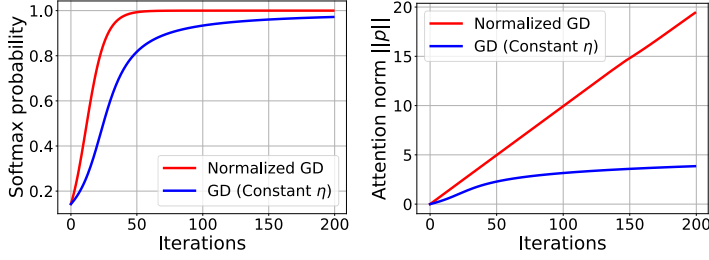
$$(\mathbf{v}_r, \mathbf{p}_R) = \arg \min_{\|\mathbf{v}\| \leq r, \|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{v}, \mathbf{p}).$$

Suppose there are token indices $\alpha = (\alpha_i)_{i=1}^n$ for which $\|\mathbf{p}^{mm}(\alpha)\|$ exists (ATT-SVM is feasible) and Assumption C holds for some $\Gamma, \nu > 0$. Then, $\lim_{R \rightarrow \infty} \mathbf{p}_R/R = \mathbf{p}^{mm}/\|\mathbf{p}^{mm}\|$, where $\mathbf{p}^{mm}$ is the solution of (ATT-SVM); and $\lim_{r \rightarrow \infty} \mathbf{v}_r/r = \mathbf{v}^{mm}/\|\mathbf{v}^{mm}\|$, where $\mathbf{v}^{mm}$ is the solution of (SVM) with $\mathbf{r}_i = \mathbf{x}_{i\alpha_i}$.

As further discussion, consider Figure 3(a) where we set $n = 3, T = d = 2$ and $\mathbf{W} = \text{Identity}$. All three inputs share the point (0, 0) which corresponds to their non-optimal tokens. The optimal tokens (denoted by $\star$) are all support vectors of the (SVM) since $\mathbf{v}^{mm} = [0, 1]$ is the optimal classifier direction (depicted by ---). Because of this, $\mathbf{p}^{mm}$ will separate optimal $\star$ tokens from tokens at the (0, 0) coordinate via (ATT-SVM) and its direction is dictated by yellow and teal colored $\star$ s which are the support vectors.

3.2 General solution when selecting one token per input

Can we relax Assumption C, and if so, what is the resulting behavior? Consider the scenario where the optimal $\mathbf{p}$ diverges to ∞ and ends up selecting one token per input. Suppose this $\mathbf{p}$ selects some coordinates $\alpha = (\alpha_i)_{i=1}^n$. Let $\mathcal{S} \subset [n]$ be the set of indices where the associated token $\mathbf{x}_{i\alpha_i}$ is a support vector when solving (SVM). Set $\bar{\mathcal{S}} = [n] - \mathcal{S}$. Our intuition is as follows: Even if we slightly perturb this $\mathbf{p}$ choice and mix other tokens $t \neq \alpha_i$ over the input set $\bar{\mathcal{S}} \subset [n]$, since $\bar{\mathcal{S}}$ is not support vector for (SVM), we can preserve the label margin (by only preserving the support vectors $\mathcal{S}$). This means that $\mathbf{p}$ may not have to enforce *max-margin* constraint over inputs $i \in \bar{\mathcal{S}}$, instead, it suffices to just select



(a) Evolution of softmax probability (b) Evolution of attention weights

Figure 4: Evolution of softmax probability and attention weights when training with normalized gradient descent or constant step size η respectively.

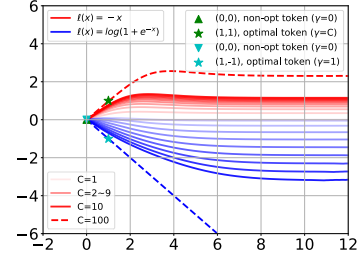


Figure 5: Trajectories of $\mathbf{p}$ with different loss functions and scores in Theorem 2.

these tokens (asymptotically). This results in the following **relaxed SVM** problem:

$$\mathbf{p}^{\text{relax}} = \arg \min_{\mathbf{p}} \|\mathbf{p}\| \quad \text{such that} \quad \mathbf{p}^\top (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}) \geq \begin{cases} 1 & \text{for all } t \neq \alpha_i, i \in \mathcal{S} \\ 0 & \text{for all } t \neq \alpha_i, i \in \bar{\mathcal{S}} \end{cases} \quad (10)$$

Here, $\mathbf{p}^\top (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}) \geq 0$ corresponds to the *selection* idea. Building on this intuition, the following theorem captures the generalized behavior of the joint regularization path.

Theorem 6 Consider the same (ERM) problem as discussed in Theorem 5. Suppose $\mathbb{S}(\mathbf{K}_i \mathbf{p}_R)_{\alpha_i} \rightarrow 1$, i.e., the tokens $(\alpha_i)_{i=1}^n$ are asymptotically selected. Let $\mathbf{v}^{\text{mm}}$ be the solution of (SVM) with $\mathbf{r}_i = \mathbf{x}_{i\alpha_i}$ and $\mathcal{S}$ be its set of support vector indices. Suppose Assumption C holds over $\mathcal{S}$ i.e. having $s_{i\alpha_i} < 1$ shrinks the margin when (SVM) is only solved over $\mathcal{S} \subset [n]$. Then, $\lim_{r \rightarrow \infty} \mathbf{v}_r / r = \mathbf{v}^{\text{mm}} / \|\mathbf{v}^{\text{mm}}\|$ and $\lim_{R \rightarrow \infty} \mathbf{p}_R / R = \mathbf{p}^{\text{relax}} / \|\mathbf{p}^{\text{relax}}\|$, where $\mathbf{p}^{\text{relax}}$ is the solution of (10) with $(\alpha_i)_{i=1}^n$ choices.

To illustrate this numerically, consider Figure 3(b) which modifies Figure 3(a) by pushing the yellow $\star$ to the northern position (0.5, 1.5). We still have $\mathbf{v}^{\text{mm}} = [0, 1]$ however the yellow $\star$ is no longer a support vector of (SVM). Thus, $\mathbf{p}$ solves the relaxed problem (10) which separates green and teal $\star$'s by enforcing the max-margin constraint on $\mathbf{p}$ (which is the red direction). Instead, yellow $\star$ only needs to achieve positive correlation with $\mathbf{p}$ (unlike Figure 3(a) where it dictates the direction). We also display the direction of $\mathbf{p}^{\text{mm}}$ using a gray dashed line.

We further investigate the evolution of softmax and logistic output probabilities throughout the training process of Figure 3(a), and the results are illustrated in Figure 3(c). The averaged softmax probability of optimal tokens is represented by the red curve and is calculated as $\frac{1}{n} \sum_{i=1}^n \max_{t \in [T]} \mathbb{S}(\mathbf{K}_i \mathbf{p})_t$. An achievement of 1 for this probability indicates that the attention mechanism successfully selects the optimal tokens. On the other hand, the logistic probability of the output is represented by the blue curve and is determined by $1/n \sum_{i=1}^n 1/(1 + e^{-Y_i f(X_i)})$. This probability also reaches a value of 1, suggesting that the inputs are correctly classified.

4 Experiments

Sparsity of softmax and evolution of attention weights. It is well known that, in practice, attention maps often exhibit sparsity and highlight salient tokens that aid inference. Our results provide a formal explanation of this when tokens are separable: Since attention selects a locally-optimal token within the input sequence and suppresses the rest, the associated attention map $\mathbb{S}(\mathbf{X}\mathbf{p})$ will (eventually) be a sparse vector. Additionally, the sparsity should arise in tandem with the increasing norm of attention weights. We provide empirical evidence to support these findings.

Synthetic experiments. Figures 4(a) and 4(b) show the evolution of the largest softmax probability and attention weights over time when using either normalized gradient or a fixed stepsize η for training. The dataset model follows Figure 1(c). The softmax probability shown in Figure 4(a) is defined as $\frac{1}{n} \sum_{i=1}^n \max_{t \in [T]} \mathbb{S}(\mathbf{K}_i \mathbf{p})_t$. When this average probability reaches the value of 1, it means attention selects only a single token per input. The attention norm in Figure 4(b), is simply equal to $\|\mathbf{p}\|$.

The red curves in both figures represent the normalized gradient method, which updates the model parameters $\mathbf{p}$ using $\mathbf{p}(t+1) = \mathbf{p}(t) - \eta \nabla \mathcal{L}(\mathbf{p}(t)) / \|\nabla \mathcal{L}(\mathbf{p}(t))\|$ with $\eta = 0.1$. The blue curves correspond

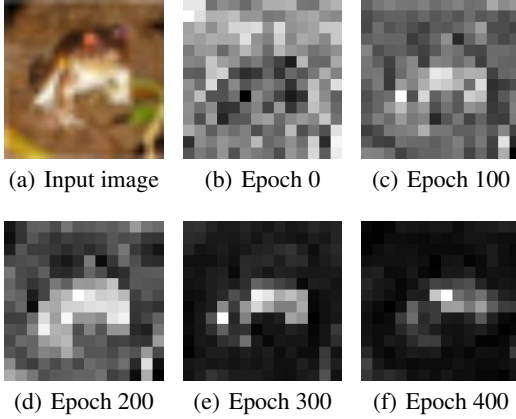


Figure 6: Illustration of the progressive change in attention weights of the [CLS] token during training in the transformer model, using a specific input image shown in Figure 6(a).

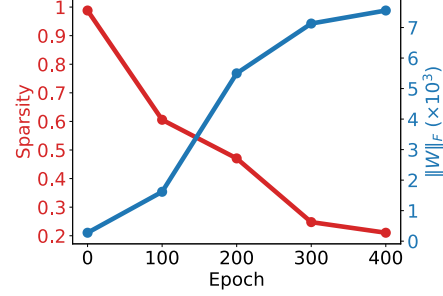


Figure 7: Red curve is the sparsity level $\widehat{\text{nnz}}(s)/T$ of the average attention map which takes values on $[0,1]$. A sparser vector implies that few key tokens receive significantly higher attention, while the majority of the tokens receive minimal attention. Blue curve is the Frobenius norm of attention weights $\|W\|_F$ of the final layer. We display their evolutions over epochs.

to gradient descent with constant learning rate given by $p(t+1) = p(t) - \eta \nabla \mathcal{L}(p(t))$ with $\eta = 1$. Observe that the normalized gradient method achieves a softmax probability of 1 quicker as vanilla GD suffers from vanishing gradients. This is visible in Figure 4(b) where blue norm curve levels off.

Real experiments. To study softmax sparsity and the evolution of attention weights throughout training, we train a vision transformer (ViT-base) model [23] from scratch, utilizing the CIFAR-10 dataset [24] for 400 epochs with fixed learning rate 3×10^{-3} . ViT tokenizes an image into 16×16 patches, thus, its softmax attention maps can be easily visualized. We examine the average attention map – associated with the [CLS] token – computed from all 12 attention heads within the model. Figure 6 provides a visual representation of the resulting attention weights (16×16 grids) corresponding to the original patch locations within the image. During the initial epochs of training, the attention weights are randomly distributed and exhibit a dense pattern. However, as the training progresses, the attention map gradually becomes sparser and the attention mechanism begins to concentrate on fewer salient patches within the image that possess distinct features that aid classification. This illustrates the evolution of attention from a random initial state to a more focused and sparse representation. These salient patches highlighted by attention conceptually corresponds to the optimal tokens within our theory.

We quantify the sparsity of the attention map via a soft-sparsity measure, denoted by $\widehat{\text{nnz}}(s)$ where s is the softmax probability vector. The soft-sparsity is computed as the ratio of the ℓ_1 -norm to the squared ℓ_2 -norm, defined as $\widehat{\text{nnz}}(s) = \|s\|_1 / \|s\|^2$. $\widehat{\text{nnz}}(s)$ takes values between 1 to $T = 256$ and a smaller value indicates a sparser vector. Also note that $\|s\|_1 = \sum_{i=1}^T s_i = 1$. Together with sparsity, Figure 7 also displays the Frobenius norm of the combined key-query matrix W of the last attention layer over epochs. The theory suggests that the increase in sparsity is associated with the growth of attention weights – which converge directionally. The results in Figure 7 align with the theory, demonstrating the progressive sparsification of the attention map as $\|W\|_F$ grows.

Transient optimization dynamics and the influence of the loss function. Theorem 2 shows that the asymptotic direction of gradient descent is determined by p^{mm*} . However, it is worth noting that transient dynamics can exhibit bias towards certain input examples and their associated optimal tokens. We illustrate this idea in Fig 5(a), which displays the trajectories of the gradients for different scores and loss functions. We consider two optimal tokens ($\star$) with scores $\gamma_1 = 1$ and $\gamma_2 = C$, where C varies. For our analysis, we examine the correlation loss $\ell(x) = -x$ and the logistic loss $\ell(x) = \log(1 + e^{-x})$.

In essence, as C increases, we can observe that the correlation loss $\ell(x) = -x$ exhibits a bias towards the token with a high score, while the logistic loss is biased towards the token with a low score. The underlying reason for this behavior can be observed from the gradients of individual inputs: $\nabla \mathcal{L}_i(p) = \ell'_i \cdot K_i^T \mathbb{S}'(Xp)Xv$, where $\mathbb{S}'(\cdot)$ represents the derivative of the softmax function and $\ell'_i := \ell'(Y_i \cdot v^T X_i^T \mathbb{S}(X_i p))$. Assuming that p (approximately) selects the optimal tokens, this

simplifies to $\ell'_i \approx \ell'(\gamma_i)$ and $\|\nabla \mathcal{L}_i(\mathbf{p})\| \propto |\ell'(\gamma_i)| \cdot \gamma_i$. With the correlation loss, $|\ell'| = 1$, resulting in $\|\nabla \mathcal{L}_i(\mathbf{p})\| \propto \gamma_i$, meaning that a larger score induces a larger gradient. On the other hand, the logistic loss behaves similarly to the exponential loss under separable data, i.e., $|\ell'| = e^{-x}/(1 + e^{-x}) \approx e^{-x}$. Consequently, $\|\nabla \mathcal{L}_i(\mathbf{p})\| \propto \gamma_i e^{-\gamma_i} \approx e^{-\gamma_i}$, indicating that a smaller score leads to a larger gradient. These observations explain the empirical behavior we observe.

5 Related Work

Implicit Regularization. The implicit bias of gradient descent in classification tasks involving separable data has been extensively examined by [22, 25, 26, 27, 28, 29]. These works typically use logistic loss or, more generally, exponentially-tailed losses to make connections to margin maximization. These results are also extended to non-separable data by [30, 31, 21]. Furthermore, there have been notable investigations into the implicit bias in regression problems/losses utilizing techniques such as mirror descent [32, 25, 33, 34, 35, 36]. In addition, several papers have explored the implicit bias of stochastic gradient descent [37, 38, 39, 40, 41, 42], as well as adaptive and momentum-based methods [43, 44, 45, 46]. Although there are similarities between our optimization approach for $\mathbf{v}$ and existing works, the optimization of $\mathbf{p}$ stands out as significantly different. Firstly, our optimization problem is nonconvex, introducing new challenges and complexities. Secondly, it necessitates the introduction of novel concepts such as locally-optimal tokens and requires a fresh analysis specifically tailored to the cones surrounding them.

Attention Mechanism. Transformers, introduced by [6], revolutionized the field of NLP and machine translation, with earlier works on self-attention by [47, 48, 49, 50]. Self-attention differs from traditional models like MLPs and CNNs by leveraging global interactions for feature representations, showing exceptional empirical performance. However, the underlying mechanisms and learning processes of the attention layer remain unknown. Recent studies such as [51, 52, 53, 54, 23] have focused on specific aspects like representing sparse functions, convex-relaxations, and expressive power. In contrast to our nonconvex (ERM), [52] studies self-attention with linear activation instead of softmax, while [53] approximates softmax using a linear operation with unit simplex constraints. Their main objective is to derive convex reformulations for ERM-based training problem. [55, 56] have developed initial results to characterize the optimization and generalization dynamics of attention. [17] is another closely related work where the authors analyze the same attention model (ERM) as us. Specifically, they jointly optimize $\mathbf{v}, \mathbf{p}$ for three gradient iterations for a contextual dataset model. However, all of these works make stringent assumptions on the data, namely, tokens are tightly clusterable or can be clearly split into clear relevant and irrelevant sets. Additionally [56] requires assumptions on initialization and [55] considers a simplified attention structure where the attention matrix is not directly parameterized with respect to the input. Our work links attention models to hard-margin SVM problems and pioneers the study of gradient descent’s implicit bias in these models.

6 Discussion

We have provided a thorough optimization-theoretic characterization of the fundamental attention model $f(\mathbf{X}) = \mathbf{v}^\top \mathbf{X}^\top \mathbb{S}(\mathbf{X}\mathbf{W}\mathbf{p})$ by formally connecting it to max-margin problems. We first established the convergence of gradient descent on $\mathbf{p}$ (or equivalently $\mathbf{W}$) in isolation. We also explored joint convergence of $(\mathbf{v}, \mathbf{p})$ via regularization path which revealed surprising implicit biases such as (10). These findings motivate several exciting avenues for future research. An immediate open problem is characterizing the (local) convergence of gradient descent for joint optimization of $(\mathbf{v}, \mathbf{p})$. Another major direction is to extend similar analysis to study self-attention layer (4) or to allow for multiple tunable tokens (where $\mathbf{p}$ becomes a matrix). Either setting will enrich the problem by allowing the attention to discover multiple hyperplanes to separate tokens. While our convergence guarantees apply when tokens are separable, it would be interesting to characterize the non-separable geometry by leveraging results developed for logistic regression analysis [31, 22]. Ideas from such earlier results can also be useful for characterizing the non-asymptotic/transient dynamics of how gradient descent aligns with the max-margin direction. Overall, we believe that max-margin token selection is a fundamental characteristic of attention mechanism and the theory developed in this work lays the groundwork of these future extensions.

Acknowledgements

This work was supported by the NSF grants CCF-2046816 and CCF-2212426, Google Research Scholar award, and Army Research Office grant W911NF2110312.

The authors express their gratitude for the valuable feedback provided by the anonymous reviewers and Christos Thrampoulidis, which has significantly improved this paper.

References

- [1] Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. Neural machine translation by jointly learning to align and translate. *The International Conference on Learning Representations*, 2015.
- [2] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, and et al. Language models are few-shot learners. In *Advances in neural information processing systems*, volume 33, pages 1877–1901, 2020.
- [3] Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, et al. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*, 2021.
- [4] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- [5] Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, et al. Palm: Scaling language modeling with pathways. *arXiv preprint arXiv:2204.02311*, 2022.
- [6] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in neural information processing systems*, volume 30, 2017.
- [7] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics.
- [8] OpenAI. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [9] Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*, 2021.
- [10] Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and Ilya Sutskever. Zero-shot text-to-image generation. In *International Conference on Machine Learning*, pages 8821–8831. PMLR, 2021.
- [11] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021.
- [12] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR, 2021.
- [13] Danny Driess, Fei Xia, Mehdi SM Sajjadi, Corey Lynch, Aakanksha Chowdhery, Brian Ichter, Ayaan Wahid, Jonathan Tompson, Quan Vuong, Tianhe Yu, et al. Palm-e: An embodied multimodal language model. *arXiv preprint arXiv:2303.03378*, 2023.

- [14] Lili Chen, Kevin Lu, Aravind Rajeswaran, Kimin Lee, Aditya Grover, Misha Laskin, Pieter Abbeel, Aravind Srinivas, and Igor Mordatch. Decision transformer: Reinforcement learning via sequence modeling. In *Advances in Neural Information Processing Systems*, volume 34, pages 15084–15097, 2021.
- [15] Scott Reed, Konrad Zolna, Emilio Parisotto, Sergio Gomez Colmenarejo, Alexander Novikov, Gabriel Barth-Maron, Mai Gimenez, Yury Sulsky, Jackie Kay, Jost Tobias Springenberg, et al. A generalist agent. *arXiv preprint arXiv:2205.06175*, 2022.
- [16] Brian Lester, Rami Al-Rfou, and Noah Constant. The power of scale for parameter-efficient prompt tuning. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 3045–3059, 2021.
- [17] Samet Oymak, Ankit Singh Rawat, Mahdi Soltanolkotabi, and Christos Thrampoulidis. On the role of attention in prompt-tuning. In *International Conference on Machine Learning*, 2023.
- [18] Xiang Lisa Li and Percy Liang. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*, 2021.
- [19] Saharon Rosset, Ji Zhu, and Trevor Hastie. Margin maximizing loss functions. *Advances in neural information processing systems*, 16, 2003.
- [20] Arun Suggala, Adarsh Prasad, and Pradeep K Ravikumar. Connecting optimization and regularization paths. *Advances in Neural Information Processing Systems*, 31, 2018.
- [21] Ziwei Ji, Miroslav Dudík, Robert E Schapire, and Matus Telgarsky. Gradient descent follows the regularization path for general losses. In *Conference on Learning Theory*, pages 2109–2136. PMLR, 2020.
- [22] Daniel Soudry, Elad Hoffer, Mor Shpigel Nacson, Suriya Gunasekar, and Nathan Srebro. The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research*, 19(1):2822–2878, 2018.
- [23] Yihe Dong, Jean-Baptiste Cordonnier, and Andreas Loukas. Attention is not all you need: Pure attention loses rank doubly exponentially with depth. In *International Conference on Machine Learning*, pages 2793–2803. PMLR, 2021.
- [24] Alex Krizhevsky, Vinod Nair, and Geoffrey Hinton. The cifar-10 dataset. *online: <http://www.cs.toronto.edu/kriz/cifar.html>*, 55(5), 2014.
- [25] Suriya Gunasekar, Jason Lee, Daniel Soudry, and Nathan Srebro. Characterizing implicit bias in terms of optimization geometry. In *International Conference on Machine Learning*, pages 1832–1841. PMLR, 2018.
- [26] Mor Shpigel Nacson, Jason Lee, Suriya Gunasekar, Pedro Henrique Pamplona Savarese, Nathan Srebro, and Daniel Soudry. Convergence of gradient descent on separable data. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pages 3420–3428. PMLR, 2019.
- [27] Ziwei Ji and Matus Telgarsky. Characterizing the implicit bias via a primal-dual analysis. In *Algorithmic Learning Theory*, pages 772–804. PMLR, 2021.
- [28] Edward Moroshko, Blake E Woodworth, Suriya Gunasekar, Jason D Lee, Nati Srebro, and Daniel Soudry. Implicit bias in deep linear classification: Initialization scale vs training accuracy. *Advances in neural information processing systems*, 33:22182–22193, 2020.
- [29] Ziwei Ji and Matus Telgarsky. Directional convergence and alignment in deep learning. In H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems*, volume 33, pages 17176–17186. Curran Associates, Inc., 2020.
- [30] Ziwei Ji and Matus Telgarsky. Risk and parameter convergence of logistic regression. *arXiv preprint arXiv:1803.07300*, 2018.

- [31] Ziwei Ji and Matus Telgarsky. The implicit bias of gradient descent on nonseparable data. In *Conference on Learning Theory*, pages 1772–1798. PMLR, 2019.
- [32] Blake Woodworth, Suriya Gunasekar, Jason D Lee, Edward Moroshko, Pedro Savarese, Itay Golan, Daniel Soudry, and Nathan Srebro. Kernel and rich regimes in overparametrized models. In *Conference on Learning Theory*, pages 3635–3673. PMLR, 2020.
- [33] Chulhee Yun, Shankar Krishnan, and Hossein Mobahi. A unifying view on implicit bias in training linear neural networks. *arXiv preprint arXiv:2010.02501*, 2020.
- [34] Tomas Vaskevicius, Varun Kanade, and Patrick Rebeschini. Implicit regularization for optimal sparse recovery. *Advances in Neural Information Processing Systems*, 32:2972–2983, 2019.
- [35] Ehsan Amid and Manfred K Warmuth. Winnowing with gradient descent. In *Conference on Learning Theory*, pages 163–182. PMLR, 2020.
- [36] Ehsan Amid and Manfred KK Warmuth. Reparameterizing mirror descent as gradient descent. *Advances in Neural Information Processing Systems*, 33:8430–8439, 2020.
- [37] Yuanzhi Li, Colin Wei, and Tengyu Ma. Towards explaining the regularization effect of initial large learning rate in training neural networks. *arXiv preprint arXiv:1907.04595*, 2019.
- [38] Guy Blanc, Neha Gupta, Gregory Valiant, and Paul Valiant. Implicit regularization for deep neural networks driven by an ornstein-uhlenbeck like process. In *Conference on learning theory*, pages 483–513. PMLR, 2020.
- [39] Jeff Z HaoChen, Colin Wei, Jason D Lee, and Tengyu Ma. Shape matters: Understanding the implicit bias of the noise covariance. *arXiv preprint arXiv:2006.08680*, 2020.
- [40] Zhiyuan Li, Tianhao Wang, and Sanjeev Arora. What happens after SGD reaches zero loss? –a mathematical framework. In *International Conference on Learning Representations*, 2022.
- [41] Alex Damian, Tengyu Ma, and Jason Lee. Label noise sgd provably prefers flat global minimizers. *arXiv preprint arXiv:2106.06530*, 2021.
- [42] Difan Zou, Jingfeng Wu, Vladimir Braverman, Quanquan Gu, Dean P Foster, and Sham Kakade. The benefits of implicit regularization from sgd in least squares problems. *Advances in Neural Information Processing Systems*, 34:5456–5468, 2021.
- [43] Qian Qian and Xiaoyuan Qian. The implicit bias of adagrad on separable data. *Advances in Neural Information Processing Systems*, 32, 2019.
- [44] Bohan Wang, Qi Meng, Huishuai Zhang, Ruoyu Sun, Wei Chen, and Zhi-Ming Ma. Momentum doesn’t change the implicit bias. *arXiv preprint arXiv:2110.03891*, 2021.
- [45] Bohan Wang, Qi Meng, Wei Chen, and Tie-Yan Liu. The implicit bias for adaptive optimization algorithms on homogeneous neural networks. In *International Conference on Machine Learning*, pages 10849–10858. PMLR, 2021.
- [46] Ziwei Ji, Nathan Srebro, and Matus Telgarsky. Fast margin maximization via dual acceleration. In *International Conference on Machine Learning*, pages 4860–4869. PMLR, 2021.
- [47] Jianpeng Cheng, Li Dong, and Mirella Lapata. Long short-term memory-networks for machine reading. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 551–561, Austin, Texas, November 2016. Association for Computational Linguistics.
- [48] Ankur Parikh, Oscar Täckström, Dipanjan Das, and Jakob Uszkoreit. A decomposable attention model for natural language inference. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 2249–2255, Austin, Texas, November 2016. Association for Computational Linguistics.

- [49] Romain Paulus, Caiming Xiong, and Richard Socher. A deep reinforced model for abstractive summarization. In *International Conference on Learning Representations*, 2018.
- [50] Zhouhan Lin, Minwei Feng, Cicero Nogueira dos Santos, Mo Yu, Bing Xiang, Bowen Zhou, and Yoshua Bengio. A structured self-attentive sentence embedding. In *International Conference on Learning Representations*, 2017.
- [51] Benjamin L Edelman, Surbhi Goel, Sham Kakade, and Cyril Zhang. Inductive biases and variable creation in self-attention mechanisms. In *International Conference on Machine Learning*, pages 5793–5831. PMLR, 2022.
- [52] Arda Sahiner, Tolga Ergen, Batu Ozturkler, John Pauly, Morteza Mardani, and Mert Pilanci. Unraveling attention via convex duality: Analysis and interpretations of vision transformers. In *International Conference on Machine Learning*, pages 19050–19088. PMLR, 2022.
- [53] Tolga Ergen, Behnam Neyshabur, and Harsh Mehta. Convexifying transformers: Improving optimization and understanding of transformer networks. *arXiv:2211.11052*, 2022.
- [54] Pierre Baldi and Roman Vershynin. The quarks of attention. *arXiv:2202.08371*, 2022.
- [55] Samy Jelassi, Michael Eli Sander, and Yanzhi Li. Vision transformers provably learn spatial structure. In Alice H. Oh, Alekh Agarwal, Danielle Belgrave, and Kyunghyun Cho, editors, *Advances in Neural Information Processing Systems*, 2022.
- [56] Hongkang Li, Meng Wang, Sijia Liu, and Pin-Yu Chen. A theoretical understanding of shallow vision transformers: Learning, generalization, and sample complexity. *arXiv preprint arXiv:2302.06015*, 2023.
- [57] Simon Du, Jason Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. Gradient descent finds global minima of deep neural networks. In *International Conference on Machine Learning*, pages 1675–1685. PMLR, 2019.
- [58] Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. *arXiv preprint arXiv:1806.07572*, 2018.
- [59] Samet Oymak and Mahdi Soltanolkotabi. Toward moderate overparameterization: Global convergence guarantees for training shallow neural networks. *IEEE Journal on Selected Areas in Information Theory*, 1(1):84–105, 2020.
- [60] Zeyuan Allen-Zhu, Yanzhi Li, and Zhao Song. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pages 242–252. PMLR, 2019.
- [61] Vladimir Vapnik. *Estimation of dependences based on empirical data*. Springer Science & Business Media, 2006.
- [62] Peter Bartlett. For valid generalization the size of the weights is more important than the size of the network. *Advances in neural information processing systems*, 9, 1996.
- [63] Albert B Novikoff. On convergence proofs for perceptrons. Technical report, STANFORD RESEARCH INST MENLO PARK CA, 1963.
- [64] Peter Bartlett, Yoav Freund, Wee Sun Lee, and Robert E Schapire. Boosting the margin: A new explanation for the effectiveness of voting methods. *The annals of statistics*, 26(5):1651–1686, 1998.
- [65] Tong Zhang and Bin Yu. Boosting with early stopping: Convergence and consistency. *Annals of Statistics*, page 1538, 2005.
- [66] Matus Telgarsky. Margins, shrinkage, and boosting. In *International Conference on Machine Learning*, pages 307–315. PMLR, 2013.

- [67] Ganesh Ramachandra Kini, Orestis Paraskevas, Samet Oymak, and Christos Thrampoulidis. Label-imbalanced and group-sensitive classification under overparameterization. *Advances in Neural Information Processing Systems*, 34:18970–18983, 2021.
- [68] Mahdi Soltanolkotabi, Dominik Stöger, and Changzhi Xie. Implicit balancing and regularization: Generalization and convergence guarantees for overparameterized asymmetric matrix sensing. *arXiv:2303.14244*, 2023.
- [69] Hossein Taheri and Christos Thrampoulidis. On generalization of decentralized learning with separable data. In *International Conference on Artificial Intelligence and Statistics*, pages 4917–4945. PMLR, 2023.
- [70] Samet Oymak and Mahdi Soltanolkotabi. Overparameterized nonlinear learning: Gradient descent takes the shortest path? In *International Conference on Machine Learning*, pages 4951–4960. PMLR, 2019.
- [71] Ziwei Ji and Matus Telgarsky. Gradient descent aligns the layers of deep linear networks. *arXiv preprint arXiv:1810.02032*, 2018.
- [72] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *Advances in Neural Information Processing Systems*, 32, 2019.
- [73] Kaifeng Lyu and Jian Li. Gradient descent maximizes the margin of homogeneous neural networks. *arXiv preprint arXiv:1906.05890*, 2019.
- [74] Lenaïc Chizat and Francis Bach. Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss. In *Conference on Learning Theory*, pages 1305–1338. PMLR, 2020.
- [75] Spencer Frei, Gal Vardi, Peter L Bartlett, and Nathan Srebro. Benign overfitting in linear classifiers and leaky relu networks from kkt conditions for margin maximization. *arXiv e-prints*, pages arXiv–2303, 2023.
- [76] Gal Vardi, Ohad Shamir, and Nati Srebro. On margin maximization in linear and relu networks. *Advances in Neural Information Processing Systems*, 35:37024–37036, 2022.
- [77] Navid Azizan, Sahin Lale, and Babak Hassibi. Stochastic mirror descent on overparameterized nonlinear models. *IEEE Transactions on Neural Networks and Learning Systems*, 33(12):7717–7727, 2021.
- [78] Navid Azizan and Babak Hassibi. Stochastic gradient/mirror descent: Minimax optimality and implicit regularization. In *International Conference on Learning Representations*.
- [79] Guorui Zhou, Xiaoqiang Zhu, Chenru Song, Ying Fan, Han Zhu, Xiao Ma, Yanghui Yan, Junqi Jin, Han Li, and Kun Gai. Deep interest network for click-through rate prediction. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1059–1068, 2018.
- [80] Qiwei Chen, Huan Zhao, Wei Li, Pipei Huang, and Wenwu Ou. Behavior sequence transformer for e-commerce recommendation in alibaba. In *Proceedings of the 1st International Workshop on Deep Learning Practice for High-Dimensional Sparse Data*, pages 1–4, 2019.
- [81] Fei Sun, Jun Liu, Jian Wu, Changhua Pei, Xiao Lin, Wenwu Ou, and Peng Jiang. Bert4rec: Sequential recommendation with bidirectional encoder representations from transformer. In *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*, pages 1441–1450, 2019.
- [82] Mia Xu Chen, Orhan Firat, Ankur Bapna, Melvin Johnson, Wolfgang Macherey, George Foster, Llion Jones, Mike Schuster, Noam Shazeer, Niki Parmar, Ashish Vaswani, Jakob Uszkoreit, Lukasz Kaiser, Zhifeng Chen, Yonghui Wu, and Macduff Hughes. The best of both worlds: Combining recent advances in neural machine translation. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 76–86, Melbourne, Australia, July 2018. Association for Computational Linguistics.

- [83] Michael Janner, Qiyang Li, and Sergey Levine. Reinforcement learning as one big sequence modeling problem. In *ICML 2021 Workshop on Unsupervised Reinforcement Learning*, 2021.
- [84] Qinqing Zheng, Amy Zhang, and Aditya Grover. Online decision transformer. In *Proceedings of the 39th International Conference on Machine Learning*, 2022.
- [85] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2020.
- [86] Hugo Touvron, Matthieu Cord, Matthijs Douze, Francisco Massa, Alexandre Sablayrolles, and Hervé Jégou. Training data-efficient image transformers & distillation through attention. In *International Conference on Machine Learning*, pages 10347–10357. PMLR, 2021.
- [87] Zi-Hang Jiang, Qibin Hou, Li Yuan, Daquan Zhou, Yujun Shi, Xiaojie Jin, Anran Wang, and Jiashi Feng. All tokens matter: Token labeling for training better vision transformers. In *Advances in Neural Information Processing Systems*, volume 34, pages 18590–18602, 2021.
- [88] Hyunjik Kim, George Papamakarios, and Andriy Mnih. The lipschitz constant of self-attention. In *International Conference on Machine Learning*, pages 5562–5571. PMLR, 2021.
- [89] Jiri Hron, Yasaman Bahri, Jascha Sohl-Dickstein, and Roman Novak. Infinite attention: Nngp and ntk for deep attention networks. In *International Conference on Machine Learning*, pages 4376–4386. PMLR, 2020.
- [90] Greg Yang. Tensor programs ii: Neural tangent kernel for any architecture. *arXiv preprint arXiv:2006.14548*, 2020.
- [91] Mostafa Dehghani, Stephan Gouws, Oriol Vinyals, Jakob Uszkoreit, and Lukasz Kaiser. Universal transformers. In *International Conference on Learning Representations*, 2018.
- [92] Chulhee Yun, Srinadh Bhojanapalli, Ankit Singh Rawat, Sashank Reddi, and Sanjiv Kumar. Are transformers universal approximators of sequence-to-sequence functions? In *International Conference on Learning Representations*, 2019.
- [93] Angeliki Giannou, Shashank Rajput, Jy-yong Sohn, Kangwook Lee, Jason D Lee, and Dimitris Papailiopoulos. Looped transformers as programmable computers. *arXiv:2301.13196*, 2023.
- [94] Yoav Levine, Noam Wies, Or Sharir, Hofit Bata, and Amnon Shashua. Limits to depth efficiencies of self-attention. In *Advances in Neural Information Processing Systems*, volume 33, pages 22640–22651, 2020.
- [95] Charlie Snell, Ruiqi Zhong, Dan Klein, and Jacob Steinhardt. Approximating how single head attention learns. *arXiv preprint arXiv:2103.07601*, 2021.
- [96] Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M Dai, and Quoc V Le. Finetuned language models are zero-shot learners. *arXiv preprint arXiv:2109.01652*, 2021.
- [97] Ekin Akyürek, Dale Schuurmans, Jacob Andreas, Tengyu Ma, and Denny Zhou. What learning algorithm is in-context learning? investigations with linear models. *arXiv:2211.15661*, 2022.
- [98] Shivam Garg, Dimitris Tsipras, Percy S Liang, and Gregory Valiant. What can transformers learn in-context? a case study of simple function classes. *Advances in Neural Information Processing Systems*, 35:30583–30598, 2022.
- [99] Yingcong Li, M Emrullah Ildiz, Dimitris Papailiopoulos, and Samet Oymak. Transformers as algorithms: Generalization and stability in in-context learning. In *International Conference on Machine Learning*, 2023.

- [100] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed Chi, Quoc Le, and Denny Zhou. Chain of thought prompting elicits reasoning in large language models. *arXiv preprint arXiv:2201.11903*, 2022.
- [101] Guhao Feng, Yuntian Gu, Bohang Zhang, Haotian Ye, Di He, and Liwei Wang. Towards revealing the mystery behind chain of thought: a theoretical perspective. *arXiv preprint arXiv:2305.15408*, 2023.
- [102] Yingcong Li, Kartik Sreenivasan, Angeliki Giannou, Dimitris Papailiopoulos, and Samet Oymak. Dissecting chain-of-thought: A study on compositional in-context learning of mlps. *arXiv preprint arXiv:2305.18869*, 2023.
- [103] Yuandong Tian, Yiping Wang, Beidi Chen, and Simon Du. Scan and snap: Understanding training dynamics and token composition in 1-layer transformer. *arXiv:2305.16380*, 2023.
- [104] Tan Minh Nguyen, Tam Minh Nguyen, Nhat Ho, Andrea L Bertozzi, Richard Baraniuk, and Stanley Osher. A primal-dual framework for transformers and neural networks. In *The Eleventh International Conference on Learning Representations*, 2023.
- [105] Davoud Ataee Tarzanagh, Yingcong Li, Christos Thrampoulidis, and Samet Oymak. Transformers as support vector machines. *arXiv preprint arXiv:2308.16898*, 2023.

Roadmap. The appendix is organized as follows: Section **A** provides basic facts about the training risk. Section **B** presents the proof of local and global gradient descent and regularized path for learning $\mathbf{p} \in \mathbb{R}^d$ with a fixed $\mathbf{v} \in \mathbb{R}^d$ choice. Section **C** provides the proof of regularized path applied to the general case of joint optimization of head $\mathbf{v}$ and attention weights $\mathbf{p}$ using a logistic loss function. Section **D** presents the regularized path applied to a more general model $f(\mathbf{X}) = \psi(\mathbf{X}^\top \mathbb{S}(\mathbf{X}\mathbf{W}^\top \mathbf{p}))$ with a nonlinear head ψ . Section **E** provides implementation details. Finally, Section **F** discusses additional related work on implicit bias and self-attention.

Table of Contents

A Addendum to Section 1	18
A.1 Preliminaries on the Training Risk	18
A.2 Proof of Lemma 1	20
B Addendum to Section 2	21
B.1 Descent and Gradient Correlation Conditions	21
B.2 Proof of Theorem 1	29
B.3 Proof of Theorem 2	30
B.4 Proof of Theorem 3	30
B.5 Proof of Theorem 4	33
B.6 Proof of Lemma 2	39
C Addendum to Section 3	39
C.1 Proof of Theorem 5	39
C.2 Proof of Theorem 6	41
D Regularization Path of Attention with Nonlinear Head	43
D.1 Proof of Theorem 8	44
D.2 Application to Linearly-mixed Labels	45
E Implementation Details and Additional Experiments	46
F Addendum to Section 5	48
F.1 Related Work on Implicit Regularization	48
F.2 Related Work on Attention Mechanism	48

A Addendum to Section 1

A.1 Preliminaries on the Training Risk

By our assumption $\psi : \mathbb{R}^d \rightarrow \mathbb{R}$ and $\ell : \mathbb{R} \rightarrow \mathbb{R}$ are differentiable functions. Recall the objective

$$\mathcal{L}(\mathbf{p}, \mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell \left(Y_i \cdot \psi(X_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p})) \right) \quad (11)$$

with the generic prediction model $\psi(\mathbf{X}^\top \mathbb{S}(\mathbf{K}\mathbf{p}))$ and $\mathbf{K} = \mathbf{X}\mathbf{W}^\top$.

Here, we write down the gradients of $\mathbf{W}$ and $\mathbf{p}$ in (11) to highlight the connection. Set $\mathbf{q} := \mathbf{W}^\top \mathbf{p}$, $\mathbf{z}\{X\} := \mathbf{X}^\top \mathbb{S}(\mathbf{K}\mathbf{p})$, and $\mathbf{a}\{X\} := \mathbf{K}\mathbf{p}$. Given X and using $\mathbf{K} = \mathbf{X}\mathbf{W}^\top$, we have that

$$\nabla_{\mathbf{q}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{X}^\top \mathbb{S}'(\mathbf{a}\{X\})\mathbf{X} \cdot \psi'(\mathbf{z}\{X\}), \quad (12a)$$

$$\nabla_{\mathbf{p}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{W}\nabla_{\mathbf{q}}\psi(\mathbf{p}, \mathbf{W}), \quad (12b)$$

$$\nabla_{\mathbf{W}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{p}\nabla_{\mathbf{q}}^\top\psi(\mathbf{p}, \mathbf{W}), \quad (12c)$$

where

$$\mathbb{S}'(\mathbf{a}\{X\}) = \text{diag}(\mathbb{S}(\mathbf{a}\{X\})) - \mathbb{S}(\mathbf{a}\{X\})\mathbb{S}(\mathbf{a}\{X\})^\top \in \mathbb{R}^{T \times T}.$$

Setting $\psi(\mathbf{z}) = \mathbf{v}^\top \mathbf{z}$ for linear head, we obtain

$$\nabla_{\mathbf{q}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{X}^\top \mathbb{S}'(\mathbf{a}\{X\})\boldsymbol{\gamma}, \quad (13a)$$

$$\nabla_{\mathbf{p}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{W}\nabla_{\mathbf{q}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{K}^\top \mathbb{S}'(\mathbf{a}\{X\})\boldsymbol{\gamma}, \quad (13b)$$

$$\nabla_{\mathbf{W}}\psi(\mathbf{p}, \mathbf{W}) = \mathbf{p}\nabla_{\mathbf{q}}^\top\psi(\mathbf{p}, \mathbf{W}) = \mathbf{p}\mathbf{v}^\top \mathbf{X}^\top \mathbb{S}'(\mathbf{a}\{X\})\mathbf{X}. \quad (13c)$$

Recalling (12b) and (12c), and defining $\ell'_i := \ell'(Y_i \cdot \psi(\mathbf{z}\{X_i\})) \in \mathbb{R}$, we have that

$$\nabla_{\mathbf{p}}\mathcal{L}(\mathbf{p}, \mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot Y_i \cdot \mathbf{W}\nabla_{\mathbf{q}}\psi(\mathbf{p}, \mathbf{W}), \quad (14a)$$

$$\nabla_{\mathbf{W}}\mathcal{L}(\mathbf{p}, \mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot Y_i \cdot \mathbf{p}\nabla_{\mathbf{q}}^\top\psi(\mathbf{p}, \mathbf{W}). \quad (14b)$$

Setting $\psi(\mathbf{z}) = \mathbf{v}^\top \mathbf{z}$ for linear head and $\boldsymbol{\gamma}_i = Y_i \cdot \mathbf{X}_i \mathbf{v}$, we obtain

$$\nabla_{\mathbf{p}}\mathcal{L}(\mathbf{p}, \mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \mathbf{K}_i^\top \mathbb{S}'(\mathbf{a}\{X_i\})\boldsymbol{\gamma}_i, \quad (15a)$$

$$\nabla_{\mathbf{W}}\mathcal{L}(\mathbf{p}, \mathbf{W}) = \mathbf{p} \left(\frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \boldsymbol{\gamma}_i^\top \mathbb{S}'(\mathbf{a}\{X_i\})\mathbf{X}_i \right). \quad (15b)$$

Lemma 3 (Key Lemma) For any $\mathbf{p}, \mathbf{q} \in \mathbb{R}^d$, let $\mathbf{a} = \mathbf{K}\mathbf{q}$, $\mathbf{s} = \mathbb{S}(\mathbf{K}\mathbf{p})$, and $\boldsymbol{\gamma} = \mathbf{X}\mathbf{v}$. Set

$$\Gamma = \sup_{t, \tau \in [T]} |\boldsymbol{\gamma}_t - \boldsymbol{\gamma}_\tau| \quad \text{and} \quad A = \sup_{t \in [T]} \|\mathbf{k}_t\| \cdot \|\mathbf{q}\|.$$

We have that

$$\left| \mathbf{a}^\top \text{diag}(\mathbf{s})\boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s}\mathbf{s}^\top \boldsymbol{\gamma} - \sum_{t \geq 2}^T (\mathbf{a}_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \right| \leq 2\Gamma A (1 - s_1)^2.$$

Proof. Set $\bar{\boldsymbol{\gamma}} = \sum_{t=1}^T \boldsymbol{\gamma}_t s_t$. We have

$$\boldsymbol{\gamma}_1 - \bar{\boldsymbol{\gamma}} = \sum_{t \geq 2}^T (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) s_t, \quad \text{and} \quad |\boldsymbol{\gamma}_1 - \bar{\boldsymbol{\gamma}}| \leq \Gamma(1 - s_1).$$

Then,

$$\begin{aligned} \mathbf{a}^\top \text{diag}(\mathbf{s})\boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s}\mathbf{s}^\top \boldsymbol{\gamma} &= \sum_{t=1}^T \mathbf{a}_t \boldsymbol{\gamma}_t s_t - \sum_{t=1}^T \mathbf{a}_t s_t \sum_{t=1}^T \boldsymbol{\gamma}_t s_t \\ &= \mathbf{a}_1 s_1 (\boldsymbol{\gamma}_1 - \bar{\boldsymbol{\gamma}}) - \sum_{t \geq 2}^T \mathbf{a}_t s_t (\bar{\boldsymbol{\gamma}} - \boldsymbol{\gamma}_t). \end{aligned} \quad (16)$$

Since

$$\left| \sum_{t \geq 2}^T \mathbf{a}_t s_t (\bar{\boldsymbol{\gamma}} - \boldsymbol{\gamma}_t) - \sum_{t \geq 2}^T \mathbf{a}_t s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \right| \leq A\Gamma(1 - s_1)^2,$$

we obtain²

$$\begin{aligned}
\mathbf{a}^\top \text{diag}(\mathbf{s})\boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s} \mathbf{s}^\top \boldsymbol{\gamma} &= \mathbf{a}_1 s_1 (\boldsymbol{\gamma}_1 - \bar{\boldsymbol{\gamma}}) - \sum_{t \geq 2}^T \mathbf{a}_t s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \pm A\Gamma(1 - s_1)^2 \\
&= \mathbf{a}_1 s_1 \sum_{t \geq 2}^T (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) s_t - \sum_{t \geq 2}^T \mathbf{a}_t s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \pm A\Gamma(1 - s_1)^2 \\
&= \sum_{t \geq 2}^T (\mathbf{a}_1 s_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \pm A\Gamma(1 - s_1)^2 \\
&= \sum_{t \geq 2}^T (\mathbf{a}_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \pm 2A\Gamma(1 - s_1)^2.
\end{aligned}$$

Here, $\pm$ on the right handside uses the fact that

$$\left| \sum_{t \geq 2}^T (\mathbf{a}_1 s_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \right| \leq (1 - s_1) A\Gamma \sum_{t \geq 2}^T s_t = (1 - s_1)^2 A\Gamma.$$

■

A.2 Proof of Lemma 1

Proof. Let us prove the result for a general step size sequence $(\eta_t)_{t \geq 0}$. On the same training data $(Y_i, X_i)_{i=1}^n$, recall the objectives $\tilde{\mathcal{L}}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot \psi(X_i^\top \mathbb{S}(X_i \mathbf{p})))$ and $\mathcal{L}(\mathbf{W}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot \psi(X_i^\top \mathbb{S}(X_i \mathbf{W}^\top \mathbf{u})))$. Suppose claim is true till iteration t . For iteration $t + 1$, using $\mathbf{W}(t)^\top \mathbf{u} = \mathbf{p}(t)$, define and observe that

$$\begin{aligned}
\mathbf{S}_i &= \mathbb{S}'(X_i \mathbf{W}(t)^\top \mathbf{u}) = \mathbb{S}'(X_i \mathbf{p}(t)), \\
s_i &= \mathbb{S}(X_i \mathbf{W}(t)^\top \mathbf{u}) = \mathbb{S}(X_i \mathbf{p}(t)), \\
\mathbf{z}\{X_i\} &= X_i^\top \mathbb{S}(X_i \mathbf{p}(t)) = X_i^\top \mathbb{S}(X_i \mathbf{W}(t)^\top \mathbf{u}),
\end{aligned}$$

for all $i \in [n]$.

Thus, using (14), we have that

$$\begin{aligned}
\nabla \tilde{\mathcal{L}}(\mathbf{p}(t)) &= \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot Y_i \cdot X_i^\top \mathbf{S}_i X_i \cdot \psi'(z\{X_i\}), \\
\nabla \mathcal{L}(\mathbf{W}(t)) &= \mathbf{u} \left(\frac{1}{n} \sum_{i=1}^n \ell'_i \cdot Y_i \cdot X_i^\top \mathbf{S}_i X_i \cdot \psi'(z\{X_i\}) \right)^\top.
\end{aligned}$$

Consequently, we found that gradient is rank-1 with left singular space equal to $\mathbf{u}$, i.e.,

$$\nabla \mathcal{L}(\mathbf{W}(t)) = \mathbf{u} \nabla^\top \tilde{\mathcal{L}}(\mathbf{p}(t)).$$

Since $\mathbf{W}(t)$'s left singular space is guaranteed to be in $\mathbf{u}$ (including $\mathbf{W}(0)$ by initialization), we only need to study the right singular vector. Using the induction till t , this yields

$$\begin{aligned}
\mathbf{W}(t+1)^\top \mathbf{u} &= \mathbf{W}(t)^\top \mathbf{u} - \eta_t \|\mathbf{u}\|^{-2} \nabla^\top \mathcal{L}(\mathbf{W}(t)) \mathbf{u} \\
&= \mathbf{p}(t) - \eta_t \|\mathbf{u}\|^{-2} \mathbf{u}^\top \mathbf{u} \nabla \tilde{\mathcal{L}}(\mathbf{p}(t)) \\
&= \mathbf{p}(t+1).
\end{aligned}$$

This concludes the induction. ■

²For simplicity, we use $\pm$ on the right hand side to denote the upper and lower bounds.

B Addendum to Section 2

B.1 Descent and Gradient Correlation Conditions

The lemma below identifies conditions under which $\mathbf{p}^{mm\star}$ is a global descent direction for $\mathcal{L}(\mathbf{p})$.

Lemma 4 *Suppose $\ell(\cdot)$ is a strictly decreasing differentiable loss function and Assumption B holds. Then, for all $\mathbf{p} \in \mathbb{R}^d$, the training loss (ERM) obeys $\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm\star} \rangle < 0$.*

Proof. Set

$$\gamma_i = Y_i \cdot X_i \mathbf{v}, \quad \mathbf{a}_i = \mathbf{K}_i \mathbf{p}, \quad \bar{\mathbf{a}}_i = \mathbf{K}_i \mathbf{p}^{mm\star}, \quad \text{and} \quad \ell'_i = \ell'(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p})). \quad (17)$$

Let us recall the gradient evaluated at $\mathbf{p}$ which is given by

$$\nabla \mathcal{L}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \mathbf{K}_i^\top \mathbb{S}'(\mathbf{a}_i) \gamma_i. \quad (18)$$

This implies that

$$\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm\star} \rangle = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \langle \bar{\mathbf{a}}_i, \mathbb{S}'(\mathbf{a}_i) \gamma_i \rangle. \quad (19)$$

To proceed, we will prove that individual summands are all strictly negative. To show that, without losing generality, let us focus on the first input and drop the subscript i for cleaner notation. This yields

$$\langle \bar{\mathbf{a}}, \mathbb{S}'(\mathbf{a}) \gamma \rangle = \bar{\mathbf{a}}^\top \text{diag}(\mathbb{S}(\mathbf{a})) \gamma - \bar{\mathbf{a}}^\top \mathbb{S}(\mathbf{a}) \mathbb{S}(\mathbf{a})^\top \gamma. \quad (20)$$

Without losing generality, assume optimal token is the first one and γ_t is a constant for all $t \geq 2$.

To proceed, we will prove the following: Suppose $\gamma = \gamma_{t \geq 2}$ is constant, $\gamma_1, \bar{\mathbf{a}}_1$ are the largest indices of $\gamma, \bar{\mathbf{a}}$. Then, for any s obeying $\sum_{t \in [T]} s_t = 1, s_t \geq 0$, we have that $\bar{\mathbf{a}}^\top \text{diag}(s) \gamma - \bar{\mathbf{a}}^\top s s^\top \gamma > 0$. To see this, we write

$$\begin{aligned} \bar{\mathbf{a}}^\top \text{diag}(s) \gamma - \bar{\mathbf{a}}^\top s s^\top \gamma &= \sum_{t=1}^T \bar{\mathbf{a}}_t \gamma_t s_t - \sum_{t=1}^T \bar{\mathbf{a}}_t s_t \sum_{t=1}^T \gamma_t s_t \\ &= \left(\bar{\mathbf{a}}_1 \gamma_1 s_1 + \gamma \sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t \right) - \left(\gamma_1 s_1 + \gamma(1 - s_1) \right) \left(\bar{\mathbf{a}}_1 s_1 + \sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t \right) \\ &= \bar{\mathbf{a}}_1 (\gamma_1 - \gamma) s_1 (1 - s_1) + \left(\gamma - (\gamma_1 s_1 + \gamma(1 - s_1)) \right) \sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t \\ &= \bar{\mathbf{a}}_1 (\gamma_1 - \gamma) s_1 (1 - s_1) - (\gamma_1 - \gamma) s_1 \sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t \\ &= (\gamma_1 - \gamma) (1 - s_1) s_1 \left[\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \right]. \end{aligned} \quad (21)$$

To proceed, let $\gamma_{\text{gap}} = \gamma_1 - \gamma$ and $\mathbf{a}_{\text{gap}} = \bar{\mathbf{a}}_1 - \max_{t \geq 2} \mathbf{a}_t$. With these, we obtain

$$\bar{\mathbf{a}}^\top \text{diag}(s) \gamma - \bar{\mathbf{a}}^\top s s^\top \gamma \geq \mathbf{a}_{\text{gap}} \gamma_{\text{gap}} s_1 (1 - s_1). \quad (22)$$

Note that

$$\begin{aligned} \mathbf{a}_{\text{gap}}^i &\geq \inf_{t \neq \text{opt}_i} (\mathbf{k}_{i \text{opt}_i} - \mathbf{k}_{it})^\top \mathbf{p}^{mm\star} \geq 1, \\ \gamma_{\text{gap}}^i &= \inf_{t \neq \text{opt}_i} \gamma_{i \text{opt}_i} - \gamma_{it} > 0, \\ s_{i1} (1 - s_{i1}) &> 0. \end{aligned}$$

On the other hand, by our assumption $\ell'_i < 0$. Hence, infimum'ing (22) over all inputs, multiplying by ℓ'_i and using (19) give the desired result. $\blacksquare$

Lemma 5 (Gradient Correlation Conditions) Consider $n = 1$ and let $\mathbf{p}^{mm} = \mathbf{p}^{mm*}$ be (ATT-SVM) solution separating $\alpha = \text{opt}$ from remaining tokens of input $\mathbf{X}$. Suppose $\ell(\cdot)$ is a strictly decreasing differentiable loss function and Assumption B holds. For any choice of $\pi > 0$, there exists $R := R_\pi$ such that, for any $\mathbf{p}$ with $\|\mathbf{p}\| \geq R$, we have

$$\left\langle \nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}}{\|\mathbf{p}\|} \right\rangle \geq (1 + \pi) \left\langle \nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle.$$

Above, observe that as $R \rightarrow \infty$, we eventually get to set $\pi = 0$.

Proof. The proof is similar to Lemma 4 at a high-level. However, we also need to account for the impact of $\mathbf{p}$ besides $\mathbf{p}^{mm}$ in the gradient correlation. The main goal is showing that $\mathbf{p}^{mm}$ is the near-optimal descent direction, thus, $\mathbf{p}$ cannot significantly outperform it.

To proceed, let $\bar{\mathbf{p}} = \|\mathbf{p}^{mm}\| \mathbf{p} / \|\mathbf{p}\|$, $M = \sup_t \|\mathbf{k}_t\|$, $\Theta = 1 / \|\mathbf{p}^{mm}\|$, $\mathbf{s} = \mathbb{S}(\mathbf{K}\mathbf{p})$, $\mathbf{a} = \mathbf{K}\bar{\mathbf{p}}$, $\bar{\mathbf{a}} = \mathbf{K}\mathbf{p}^{mm}$. Without losing generality assume $\text{opt} = 1$. Set $\gamma = \gamma_{t \geq 2}$. Repeating the proof of Lemma 4 yields

$$\begin{aligned} \langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm} \rangle &= \ell' \cdot (\gamma_1 - \gamma)(1 - s_1)s_1 \left[\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \right], \\ \langle \nabla \mathcal{L}(\mathbf{p}), \bar{\mathbf{p}} \rangle &= \ell' \cdot (\gamma_1 - \gamma)(1 - s_1)s_1 \left[\mathbf{a}_1 - \frac{\sum_{t \geq 2}^T \mathbf{a}_t s_t}{\sum_{t \geq 2}^T s_t} \right]. \end{aligned}$$

Given π , for sufficiently large R , we wish to show that

$$\mathbf{a}_1 - \frac{\sum_{t \geq 2}^T \mathbf{a}_t s_t}{\sum_{t \geq 2}^T s_t} \leq (1 + \pi) \cdot \left[\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \right]. \quad (23)$$

We consider two scenarios.

Scenario 1: $\|\bar{\mathbf{p}} - \mathbf{p}^{mm}\| \leq \epsilon := \pi / (2M)$. In this scenario, for any token, we find that

$$|\mathbf{a}_t - \bar{\mathbf{a}}_t| = |\mathbf{k}_t^\top (\bar{\mathbf{p}} - \mathbf{p}^{mm})| \leq M \|\bar{\mathbf{p}} - \mathbf{p}^{mm}\| \leq M\epsilon.$$

Consequently, we obtain

$$\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \geq \mathbf{a}_1 - \frac{\sum_{t \geq 2}^T \mathbf{a}_t s_t}{\sum_{t \geq 2}^T s_t} - 2M\epsilon = \mathbf{a}_1 - \frac{\sum_{t \geq 2}^T \mathbf{a}_t s_t}{\sum_{t \geq 2}^T s_t} - \pi.$$

Also noticing $\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \geq 1$ (thanks to $\mathbf{p}^{mm}$ satisfying ≥ 1 margin), this implies (23).

Scenario 2: $\|\bar{\mathbf{p}} - \mathbf{p}^{mm}\| \geq \epsilon := \pi / (2M)$. In this scenario, for some $\nu = \nu(\epsilon)$ and $\tau \geq 2$, we have that

$$\bar{\mathbf{p}}^\top (\mathbf{k}_1 - \mathbf{k}_\tau) = \mathbf{a}_1 - \mathbf{a}_\tau \leq 1 - 2\nu.$$

Here $\tau = \arg \max_{t \geq 2} \bar{\mathbf{p}}^\top \mathbf{k}_t$ denotes the nearest point to $\mathbf{k}_1$. Recall that $\mathbf{s} = \mathbb{S}(\bar{\mathbf{R}}\mathbf{a})$ where $\bar{\mathbf{R}} = \|\mathbf{p}\| / \|\mathbf{p}^{mm}\|$. To proceed, split the tokens into two groups: Let $\mathcal{N}$ be the group of tokens obeying $\bar{\mathbf{p}}^\top (\mathbf{k}_1 - \mathbf{k}_t) \leq 1 - \nu$ for $t \in \mathcal{N}$ and $[T] - \mathcal{N}$ be the rest. Observe that

$$\frac{\sum_{t \in [T] - \mathcal{N}} s_t}{\sum_{t \geq 2}^T s_t} \leq \frac{\sum_{t \in [T] - \mathcal{N}} s_t}{s_\tau} \leq T \frac{e^{\nu \bar{\mathbf{R}}}}{e^{2\nu \bar{\mathbf{R}}}} = T e^{-\bar{\mathbf{R}}\nu}.$$

Set $\bar{M} = M/\Theta$ and note that $\|\mathbf{a}_t\| \leq \|\mathbf{p}^{mm}\| \cdot \|\mathbf{k}_t\| \leq \bar{M}$. Using $\bar{\mathbf{p}}^\top (\mathbf{k}_1 - \mathbf{k}_t) \leq 1 - \nu$ over $t \in \mathcal{N}$ and plugging in the above bound, we obtain

$$\begin{aligned} \frac{\sum_{t \geq 2}^T (\mathbf{a}_1 - \mathbf{a}_t)s_t}{\sum_{t \geq 2}^T s_t} &= \frac{\sum_{t \in \mathcal{N}} (\mathbf{a}_1 - \mathbf{a}_t)s_t}{\sum_{t \geq 2}^T s_t} + \frac{\sum_{t \in [T] - \mathcal{N}} (\mathbf{a}_1 - \mathbf{a}_t)s_t}{\sum_{t \geq 2}^T s_t} \\ &\leq 1 - \nu + 2\bar{M}T e^{-\bar{\mathbf{R}}\nu}. \end{aligned}$$

Using the fact that $\bar{\mathbf{a}}_1 - \frac{\sum_{t \geq 2}^T \bar{\mathbf{a}}_t s_t}{\sum_{t \geq 2}^T s_t} \geq 1$, the above implies (23) with $\pi' = 2\bar{M}T e^{-\bar{\mathbf{R}}\nu} - \nu$. To proceed, choose $R_\pi = \nu^{-1} \Theta^{-1} \log(2\bar{M}T/\pi)$ to ensure $\pi' \leq \pi$. ■

The following lemma states the descent property of gradient descent for $\mathcal{L}(\mathbf{p})$ under Assumption A. It is important to note that although the infimum of the optimization problem is $\mathcal{L}^*$, it is not achieved at any finite $\mathbf{p}$. Additionally, there are no finite critical points $\mathbf{p}$.

Lemma 6 Under Assumption A, the function $\mathcal{L}(\mathbf{p})$ is L_p -smooth, where

$$L_p := \frac{1}{n} \sum_{i=1}^n \left(M_0 \|\mathbf{v}\|^2 \|\mathbf{W}\|^2 \|\mathbf{X}_i\|^4 + 3M_1 \|\mathbf{v}\| \|\mathbf{W}\|^2 \|\mathbf{X}_i\|^3 \right). \quad (24)$$

Furthermore, if $\eta \leq 1/L_p$, then, for any initialization $\mathbf{p}(0)$, with the GD sequence $\mathbf{p}(t+1) = \mathbf{p}(t) - \eta \nabla \mathcal{L}(\mathbf{p}(t))$, we have

$$\mathcal{L}(\mathbf{p}(t+1)) - \mathcal{L}(\mathbf{p}(t)) \leq -\frac{\eta}{2} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2, \quad (25)$$

for all $t \geq 0$. This implies that

$$\sum_{t=0}^{\infty} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 < \infty, \text{ and } \lim_{t \rightarrow \infty} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 = 0. \quad (26)$$

Proof. Recall that we defined $\gamma_i = Y_i \cdot \mathbf{X}_i \mathbf{v}$ and $\mathbf{a}_i = \mathbf{K}_i \mathbf{p}$. The gradient of $\mathcal{L}(\mathbf{p})$ is given by

$$\nabla \mathcal{L}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}) \right) \cdot \mathbf{K}_i^\top \mathbb{S}'(\mathbf{a}_i) \gamma_i.$$

Note that for any $\mathbf{p} \in \mathbb{R}^d$, the Jacobian of $\mathbb{S}(\mathbf{K}_i \mathbf{p})$ is given by

$$\frac{\partial \mathbb{S}(\mathbf{K}_i \mathbf{p})}{\partial \mathbf{p}} = \mathbb{S}'(\mathbf{K}_i \mathbf{p}) \mathbf{K}_i = \left(\text{diag}(\mathbb{S}(\mathbf{K}_i \mathbf{p})) - \mathbb{S}(\mathbf{K}_i \mathbf{p}) \mathbb{S}(\mathbf{K}_i \mathbf{p})^\top \right) \mathbf{K}_i. \quad (27)$$

The Jacobian (27) together with the definition of the softmax function $\mathbb{S}(\cdot)$ implies that $\|\partial \mathbb{S}(\mathbf{K}_i \mathbf{p}) / \partial \mathbf{p}\| \leq \|\mathbf{K}_i\|$. Hence, for any $\mathbf{p}, \dot{\mathbf{p}} \in \mathbb{R}^d$, we have

$$\|\mathbb{S}(\mathbf{K}_i \mathbf{p}) - \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}})\| \leq \|\mathbf{K}_i\| \|\mathbf{p} - \dot{\mathbf{p}}\|, \quad (28a)$$

and

$$\begin{aligned} \|\mathbb{S}'(\mathbf{K}_i \mathbf{p}) - \mathbb{S}'(\mathbf{K}_i \dot{\mathbf{p}})\| &\leq \|\text{diag}(\mathbb{S}(\mathbf{K}_i \mathbf{p})) - \text{diag}(\mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}}))\| \\ &\quad + \|\mathbb{S}(\mathbf{K}_i \mathbf{p}) \mathbb{S}(\mathbf{K}_i \mathbf{p})^\top - \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}}) \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}})^\top\| \\ &\leq 3\|\mathbf{K}_i\| \|\mathbf{p} - \dot{\mathbf{p}}\|. \end{aligned} \quad (28b)$$

Here, the last inequality uses the fact that $|ab - cd| \leq |d||a - c| + |a||b - d|$.

Next, for any $\mathbf{p}, \dot{\mathbf{p}} \in \mathbb{R}^d$, we have

$$\begin{aligned} \|\nabla \mathcal{L}(\mathbf{p}) - \nabla \mathcal{L}(\dot{\mathbf{p}})\| &\leq \frac{1}{n} \sum_{i=1}^n \left\| \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}) \right) \cdot \mathbf{K}_i^\top \mathbb{S}'(\mathbf{K}_i \mathbf{p}) \gamma_i - \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}}) \right) \cdot \mathbf{K}_i^\top \mathbb{S}'(\mathbf{K}_i \dot{\mathbf{p}}) \gamma_i \right\| \\ &\leq \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{K}_i^\top \mathbb{S}'(\mathbf{K}_i \dot{\mathbf{p}}) \gamma_i \right\| \left| \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}) \right) - \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}}) \right) \right| \\ &\quad + \frac{1}{n} \sum_{i=1}^n \left| \ell' \left(\gamma_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}) \right) \right| \left\| \mathbf{K}_i^\top \mathbb{S}'(\mathbf{K}_i \mathbf{p}) \gamma_i - \mathbf{K}_i^\top \mathbb{S}'(\mathbf{K}_i \dot{\mathbf{p}}) \gamma_i \right\| \\ &\leq \frac{1}{n} \sum_{i=1}^n M_0 \|\gamma_i\|^2 \|\mathbf{K}_i\| \|\mathbb{S}(\mathbf{K}_i \mathbf{p}) - \mathbb{S}(\mathbf{K}_i \dot{\mathbf{p}})\| \\ &\quad + \frac{1}{n} \sum_{i=1}^n M_1 \|\gamma_i\| \|\mathbf{K}_i\| \|\mathbb{S}'(\mathbf{K}_i \mathbf{p}) - \mathbb{S}'(\mathbf{K}_i \dot{\mathbf{p}})\|, \end{aligned} \quad (29)$$

where the second inequality follows from the fact that $|ab - cd| \leq |d||a - c| + |a||b - d|$ and the third inequality uses Assumption A.

Substituting (28a) and (28b) into (29), we get

$$\begin{aligned} \|\nabla \mathcal{L}(\mathbf{p}) - \nabla \mathcal{L}(\dot{\mathbf{p}})\| &\leq \frac{1}{n} \sum_{i=1}^n \left(M_0 \|\gamma_i\|^2 \|\mathbf{K}_i\|^2 + 3M_1 \|\mathbf{K}_i\|^2 \|\gamma_i\| \right) \|\mathbf{p} - \dot{\mathbf{p}}\| \\ &\leq \frac{1}{n} \sum_{i=1}^n \left(M_0 \|\mathbf{v}\|^2 \|\mathbf{W}\|^2 \|\mathbf{X}_i\|^4 + 3M_1 \|\mathbf{v}\| \|\mathbf{W}\|^2 \|\mathbf{X}_i\|^3 \right) \|\mathbf{p} - \dot{\mathbf{p}}\| \\ &\leq L_p \|\mathbf{p} - \dot{\mathbf{p}}\|, \end{aligned}$$

where L_p is defined in (24).

The remaining proof follows standard gradient descent analysis (see e.g. [22, Lemma 10]). Since $\mathcal{L}(\mathbf{p})$ is L_p -smooth, we get

$$\begin{aligned}\mathcal{L}(\mathbf{p}(t+1)) &\leq \mathcal{L}(\mathbf{p}(t)) + \nabla \mathcal{L}(\mathbf{p}(t))^\top (\mathbf{p}(t+1) - \mathbf{p}(t)) + \frac{L_p}{2} \|\mathbf{p}(t+1) - \mathbf{p}(t)\|^2 \\ &= \mathcal{L}(\mathbf{p}(t)) - \eta \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 + \frac{L_p \eta^2}{2} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 \\ &= \mathcal{L}(\mathbf{p}(t)) - \eta \left(1 - \frac{L_p \eta}{2}\right) \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 \\ &\leq \mathcal{L}(\mathbf{p}(t)) - \frac{\eta}{2} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2,\end{aligned}$$

where the last inequality follows from our assumption on the stepsize.

The above inequality implies that

$$\sum_{t=0}^{\infty} \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2 \leq \frac{2}{\eta} (\mathcal{L}(\mathbf{p}(0)) - \mathcal{L}^*), \quad (30)$$

where the right hand side is upper bounded by a finite constant. This is because, by Assumption A, $\mathcal{L}(\mathbf{p}(0)) < \infty$ and $\mathcal{L}^* \leq \mathcal{L}(\mathbf{p}(t))$, where $\mathcal{L}^*$ denotes the minimum objective.

Finally, (30) yields the expression (26). ■

In the following lemma, we demonstrate the existence of parameters $\mu = \mu(\alpha) > 0$ and $R_\mu > 0$ such that when R_μ is sufficiently large, there are no stationary points within $C_{\mu, R_\mu}(\mathbf{p}^{mm})$. Additionally, we provide the local gradient correlation condition.

Lemma 7 (Local Gradient Condition) *Suppose Assumption A on the loss function ℓ holds. Let $\alpha = (\alpha_i)_{i=1}^n$ be indices of locally-optimal tokens per Definition 2.*

L1. *There exists a positive scalar $\mu = \mu(\alpha) > 0$ such that for sufficiently large $\bar{R}_\mu$, no stationary point exists within $C_{\mu, \bar{R}_\mu}(\mathbf{p}^{mm})$, where $C_{\mu, \bar{R}_\mu}$ is defined in (8).*

L2. *For all $\mathbf{q}, \mathbf{p} \in \text{cone}_\mu(\mathbf{p}^{mm})$ with $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$ and $\|\mathbf{p}\| \geq \bar{R}_\mu$ with same $\bar{R}_\mu$ choice as (L1), there exist dataset dependent constants $C, c > 0$ such that*

$$C \cdot \frac{1}{n} \sum_{i \in [n]} \{1 - \mathbb{S}(\mathbf{K}_i \mathbf{p})_{\alpha_i}\} \geq -\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle \geq c \cdot \frac{1}{n} \sum_{i \in [n]} \{1 - \mathbb{S}(\mathbf{K}_i \mathbf{p})_{\alpha_i}\} > 0, \quad (31a)$$

$$\|\nabla \mathcal{L}(\mathbf{p})\| \leq \bar{A} C \cdot \frac{1}{n} \sum_{i \in [n]} \{1 - \mathbb{S}(\mathbf{K}_i \mathbf{p})_{\alpha_i}\} \leq \bar{A} C T e^{-\bar{R}_\mu \Theta / 2}. \quad (31b)$$

$$-\left\langle \frac{\mathbf{q}}{\|\mathbf{q}\|}, \frac{\nabla \mathcal{L}(\mathbf{p})}{\|\nabla \mathcal{L}(\mathbf{p})\|} \right\rangle \geq \frac{c \Theta}{C \bar{A}} > 0, \quad (31c)$$

Here, $\bar{A} = \max_{i \in [n], t, \tau \in [T]} \|\mathbf{k}_{it} - \mathbf{k}_{i\tau}\|$ and $\Theta = 1/\|\mathbf{p}^{mm}\|$.

L3. *For any $\pi > 0$, there exists R_π such that $R_\pi \geq \bar{R}_\mu$ and all $\mathbf{p} \in C_{\mu, R_\pi}(\mathbf{p}^{mm})$ obeys*

$$\left\langle \nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}}{\|\mathbf{p}\|} \right\rangle \geq (1 + \pi) \left\langle \nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle.$$

Proof. Let $\mathbf{p}^{mm} = \mathbf{p}^{mm}(\alpha)$ be the solution of (ATT-SVM). Recall

$$C_{\mu, \bar{R}_\mu}(\mathbf{p}^{mm}) = \text{cone}_\mu(\mathbf{p}^{mm}) \cap \{\mathbf{p} \mid \|\mathbf{p}\| \geq \bar{R}_\mu\}.$$

Let $(\mathcal{T}_i)_{i=1}^n$ be the sets of all SVM-neighbors per Definition 2. Let $\bar{\mathcal{T}}_i = [T] - \mathcal{T}_i - \{\alpha_i\}$ be the set of non-SVM-neighbor tokens, $i \in [n]$. Let

$$\begin{aligned}\Theta &= 1/\|\mathbf{p}^{mm}\|, \\ \delta &= 0.5 \min_{i \in [n]} \min_{t \in \mathcal{T}_i, \tau \in \bar{\mathcal{T}}_i} (\mathbf{k}_{it} - \mathbf{k}_{i\tau})^\top \mathbf{p}^{mm}, \\ A &= \max_{i \in [n], t \in [T]} \|\mathbf{k}_{it}\|/\Theta, \\ \mu &\leq \mu(\delta) = \frac{1}{8} \left(\frac{\min(0.5, \delta)}{A} \right)^2.\end{aligned}\tag{32}$$

When $\bar{\mathcal{T}}_i = \emptyset$ for all $i \in [n]$ (i.e. globally-optimal indices), we set $\delta = \infty$ as all non-neighbor related terms will disappear. Since $\mathbf{p}^{mm}$ is the max-margin model ensuring $(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{p}^{mm} \geq 1$ for all $i \in [n]$, the following inequalities hold for all $\mathbf{q} \in \text{cone}_\mu(\mathbf{p}^{mm})$, $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$ and all $i \in [n]$, $t \in \mathcal{T}_i$, $\tau \in \bar{\mathcal{T}}_i$:

$$\begin{aligned}(\mathbf{k}_{it} - \mathbf{k}_{i\tau})^\top \mathbf{q} &\geq \delta > 0, \\ (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{i\tau})^\top \mathbf{q} &\geq 1 + \delta, \\ 3/2 &\geq (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{q} \geq 1/2.\end{aligned}\tag{33}$$

Here, we used $\|\mathbf{q} - \mathbf{p}^{mm}\|^2 / \|\mathbf{p}^{mm}\|^2 \leq 2\mu$ which implies $\|\mathbf{q} - \mathbf{p}^{mm}\| \leq \sqrt{2\mu}/\Theta$.

L1. and **L2.** Now that the choice of local cone is determined, we need to prove the main claims. We will lower bound $-\mathbf{q}^\top \nabla \mathcal{L}(\mathbf{p})$ and establish its strict positivity for $\|\mathbf{p}\| \geq R$, where $R = \bar{R}_\mu$. This will show that there is no stationary point as a by product.

Consider any $\mathbf{q} \in \mathbb{R}^d$ satisfying $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$. To proceed, we write the gradient correlation following (18) and (21)

$$\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle = \frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \langle \mathbf{a}_i, \mathbb{S}'(\mathbf{a}'_i) \gamma_i \rangle,\tag{34}$$

where we denoted $\ell'_i = \ell'(Y_i \cdot \mathbf{v}^\top X_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}))$, $\mathbf{a}_i = \mathbf{K}_i \mathbf{q}$, $\mathbf{a}'_i = \mathbf{K}_i \mathbf{p}$, $s_i = \mathbb{S}(\mathbf{K}_i \mathbf{p})$.

Using (33), for all $t \in \mathcal{T}_i$, $\tau \in \bar{\mathcal{T}}_i$, for all $\mathbf{p} \in C_{\mu, R}(\mathbf{p}^{mm})$, we have that

$$\mathbf{a}'_{i\alpha_i} - \mathbf{a}'_{i\tau} \geq R\Theta(1 + \delta), \quad \text{and} \quad \mathbf{a}'_{it} - \mathbf{a}'_{i\tau} \geq R\Theta\delta.$$

Consequently, we can bound the softmax probabilities $s_i = \mathbb{S}(\mathbf{K}_i \mathbf{p})$ as follows: For all $i \in [n]$,

$$\begin{aligned}S_i &:= \sum_{\tau \in \mathcal{T}_i} s_{i\tau} \leq 1 - s_{i\alpha_i} = \sum_{\tau \neq \alpha_i} s_{i\tau} \leq T e^{-R\Theta/2} s_{i\alpha_i} \leq T e^{-R\Theta/2}, \\ Q_i &:= \sum_{\tau \in \bar{\mathcal{T}}_i} s_{i\tau} \leq T e^{-R\Theta\delta} s_{i\alpha_i} \leq T e^{-R\Theta\delta} S_i, \quad \forall t_i \in \mathcal{T}_i.\end{aligned}\tag{35}$$

Recall scores $\gamma_{it} = Y_i \cdot \mathbf{v}^\top \mathbf{x}_{it}$. Define the score gaps over neighbors:

$$\gamma_i^{\text{gap}} = \gamma_{i\alpha_i} - \max_{t \in \mathcal{T}_i} \gamma_{it}, \quad \text{and} \quad \bar{\gamma}_i^{\text{gap}} = \gamma_{i\alpha_i} - \min_{t \in \mathcal{T}_i} \gamma_{it}.$$

It follows from (32) that

$$A = \max_{i \in [n], t \in [T]} \|\mathbf{k}_{it}\|/\Theta \geq \max_{i \in [n], t \in [T]} \|\mathbf{a}_{it}\| = \|\mathbf{k}_{it} \mathbf{q}\|.$$

Define the α -dependent global scalar $\Gamma = \sup_{i \in [n], t, \tau \in [T]} |\gamma_{it} - \gamma_{i\tau}|$. Let us focus on a fixed datapoint $i \in [n]$, assume (without losing generality) $\alpha_i = 1$, and drop subscripts i , that is, $\alpha := \alpha_i = 1$, $\mathbf{X} := \mathbf{X}_i$, $Y := Y_i$, $\mathbf{K} := \mathbf{K}_i$, $\mathbf{a}' = \mathbf{K} \mathbf{p}$, $\mathbf{a} = \mathbf{K} \mathbf{q}$, $s = \mathbb{S}(\mathbf{K} \mathbf{p})$, $\gamma = Y \cdot \mathbf{X} \mathbf{v}$, $\gamma^{\text{gap}} := \gamma_i^{\text{gap}}$, $\bar{\gamma}^{\text{gap}} := \bar{\gamma}_i^{\text{gap}}$, $Q := Q_i$, and $S := S_i$. Directly applying Lemma 3, we obtain

$$\left| \mathbf{a}^\top \text{diag}(s) \gamma - \mathbf{a}^\top s s^\top \gamma - \sum_{t \geq 2} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\gamma_1 - \gamma_t) \right| \leq 2\Gamma A (1 - s_1)^2.$$

To proceed, let us decouple the non-neighbors within $\sum_{t \geq 2} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\gamma_1 - \gamma_t)$ via

$$\left| \sum_{t \in \bar{\mathcal{T}}} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\gamma_1 - \gamma_t) \right| \leq 2Q\Gamma A.$$

Aggregating these, we found

$$\left| \mathbf{a}^\top \text{diag}(\mathbf{s}) \boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s} \mathbf{s}^\top \boldsymbol{\gamma} - \sum_{t \in \mathcal{T}_i} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \right| \leq 2\Gamma A ((1 - s_1)^2 + Q). \quad (36)$$

To proceed, let us upper/lower bound the gradient correlation. We use two bounds depending on $\mathbf{q} \in \text{cone}_\mu(\mathbf{p}^{mm})$ (Case 1) or general $\mathbf{q} \in \mathbb{R}^d$ (Case 2).

- Case 1: $\mathbf{q} \in \text{cone}_\mu(\mathbf{p}^{mm})$. Since $1.5 \geq \mathbf{a}_1 - \mathbf{a}_t \geq 0.5$ following (33), we find

$$1.5 \cdot S \cdot \bar{\gamma}^{gap} \geq \sum_{t \in \mathcal{T}_i} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t) \geq 0.5 \cdot S \cdot \gamma^{gap}.$$

Next we claim that S dominates $((1 - s_1)^2 + Q)$ for large R . Specifically, we wish for

$$S \cdot \gamma^{gap} / 4 \geq 4\Gamma A \max((1 - s_1)^2, Q) \iff S \geq 16 \frac{\Gamma A}{\gamma^{gap}} \max((1 - s_1)^2, Q). \quad (37)$$

Now choose $R \geq \delta^{-1} \log(T)/\Theta$ to ensure $Q \leq S$ since $Q \leq T e^{-R\Theta\delta} S$. Consequently

$$(1 - s_1)^2 = (Q + S)^2 \leq 4S^2 \leq 4ST e^{-R\Theta/2}.$$

Combining these, what we wish is ensured by guaranteeing

$$S \geq 16 \frac{\Gamma A}{\gamma^{gap}} \max(4ST e^{-R\Theta/2}, T e^{-R\Theta\delta} S). \quad (38)$$

This in turn is ensured for all inputs $i \in [n]$ by choosing

$$R = \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{64T\Gamma A}{\gamma_{\min}^{gap}} \right), \quad (39)$$

where $\gamma_{\min}^{gap} = \min_{i \in [n]} \gamma_i^{gap}$ is the global scalar which is the worst case score gap over all inputs. With the above choice of R we guaranteed

$$2(1 - s_1) \cdot \bar{\gamma}^{gap} \geq 2 \cdot S \cdot \bar{\gamma}^{gap} \geq \mathbf{a}^\top \text{diag}(\mathbf{s}) \boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s} \mathbf{s}^\top \boldsymbol{\gamma} \geq \frac{S \cdot \gamma^{gap}}{4} \geq \frac{(1 - s_1) \gamma^{gap}}{8}.$$

Since this holds over all inputs, going back to the gradient correlation (34) and averaging above over all inputs $i \in [n]$ and plugging back the indices i , we obtain the advertised bound by setting $q_i = 1 - s_{i\alpha_i}$ (where we set $\alpha_i = 1$ above without losing generality)

$$\frac{2}{n} \sum_{i \in [n]} -\ell'_i \cdot q_i \cdot \bar{\gamma}_i^{gap} \geq -\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle \geq \frac{1}{8n} \sum_{i \in [n]} -\ell'_i \cdot q_i \cdot \gamma_i^{gap}. \quad (40)$$

Let $-\ell'_{\min/\max}$ be the min/max values negative loss derivative admits over the ball $[-B, B]$ for $B = \|\mathbf{v}\| \cdot \max_{i,t} \|\mathbf{x}_{it}\|$ and note that $\max_{i \in [n]} \bar{\gamma}_i^{gap} > 0$ and $\min_{i \in [n]} \gamma_i^{gap} > 0$ are dataset dependent constants. Then, we declare the constants $C = -2\ell'_{\max} \cdot \max_{i \in [n]} \bar{\gamma}_i^{gap} > 0$, $c = -(1/8)\ell'_{\min} \cdot \min_{i \in [n]} \gamma_i^{gap} > 0$ to obtain the bound

$$\frac{C}{n} \sum_{i \in [n]} q_i \geq -\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle \geq \frac{c}{n} \sum_{i \in [n]} q_i, \quad (41)$$

which is the desired statement in (31a).

- Case 2: $\mathbf{q} \in \mathbb{R}^d$ and $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$. Define $\bar{A} = \max_{i \in [n], t, \tau \in [T]} \|\mathbf{k}_{it} - \mathbf{k}_{i\tau}\|$. For any $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$, we use the fact that

$$\|\mathbf{a}_1 - \mathbf{a}_t\| \leq \|\mathbf{k}_1 - \mathbf{k}_t\| \cdot \|\mathbf{q}\| \leq \frac{\bar{A}}{\Theta}.$$

Note that by definition $\frac{\bar{A}}{\Theta} \geq 1$. To proceed, we can upper bound

$$\frac{\bar{A}}{\Theta} \cdot S \cdot \bar{\gamma}^{gap} \geq \sum_{t \in \mathcal{T}} (\mathbf{a}_1 - \mathbf{a}_t) s_t (\boldsymbol{\gamma}_1 - \boldsymbol{\gamma}_t). \quad (42)$$

By choosing the same R as in (39) to ensure S dominates $((1 - s_1)^2 + Q)$ and since $\frac{\bar{A}}{\Theta} \geq 1$, we guaranteed

$$\frac{2\bar{A}}{\Theta} \cdot S \cdot \bar{\gamma}^{gap} \geq \mathbf{a}^\top \text{diag}(\mathbf{s}) \boldsymbol{\gamma} - \mathbf{a}^\top \mathbf{s} \mathbf{s}^\top \boldsymbol{\gamma}.$$

Going back to the gradient correlation (34) and averaging above over all inputs $i \in [n]$, with the same definition of $C > 0$, we obtain

$$\frac{\bar{A}C}{\Theta n} \sum_{i \in [n]} q_i \geq -\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle. \quad (43)$$

To proceed, since (43) holds for any $\mathbf{q} \in \mathbb{R}^d$ and $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$, we observe that when choosing $\mathbf{q} = \frac{\|\mathbf{p}^{mm}\|}{\|\nabla \mathcal{L}(\mathbf{p})\|} \cdot \nabla \mathcal{L}(\mathbf{p})$, this implies that

$$\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{q} \rangle = \|\nabla \mathcal{L}(\mathbf{p})\| \cdot \|\mathbf{p}^{mm}\| \leq \frac{\bar{A}C}{\Theta n} \sum_{i \in [n]} q_i.$$

Simplifying $\Theta = 1/\|\mathbf{p}^{mm}\|$ on both sides yields (31b). Incorporating (35) in the bound above provides the exponential upper bound that decay with R .

Combining this with (41), we obtain that for all $\mathbf{q}, \mathbf{p} \in \text{cone}_\mu(\mathbf{p}^{mm})$ and $\|\mathbf{q}\| \geq \bar{R}_\mu$

$$-\left\langle \frac{\mathbf{q}}{\|\mathbf{q}\|}, \frac{\nabla \mathcal{L}(\mathbf{p})}{\|\nabla \mathcal{L}(\mathbf{p})\|} \right\rangle \geq \frac{c\Theta}{C\bar{A}}.$$

This gives the desired result in (31c).

L3.: Establishing gradient correlation.

Our final goal is establishing gradient comparison between $\mathbf{p}, \mathbf{p}^{mm}$ for the same choice of $\mu > 0$ provided in (32). Define $\bar{\mathbf{p}} = \|\mathbf{p}^{mm}\| \mathbf{p} / \|\mathbf{p}\|$ to be the normalized vector. Set notations $\mathbf{a}_i = \mathbf{K}_i \bar{\mathbf{p}}$, $\bar{\mathbf{a}}_i = \mathbf{K}_i \mathbf{p}^{mm}$, and $\boldsymbol{\gamma}_i = Y_i \cdot \mathbf{X}_i \mathbf{v}$.

To establish the result, using (34), we will prove that, for any $\pi > 0$, there is sufficiently large $R = R_\pi$ such that for any $\mathbf{p} \in C_{\mu, R}(\mathbf{p}^{mm})$:

$$\begin{aligned} \left\langle -\nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}}{\|\mathbf{p}\|} \right\rangle &= -\frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \langle \mathbf{a}_i, \mathbb{S}'(\mathbf{K}_i \mathbf{p}) \boldsymbol{\gamma}_i \rangle \\ &\leq -\frac{1+\pi}{n} \sum_{i=1}^n \ell'_i \cdot \langle \bar{\mathbf{a}}_i, \mathbb{S}'(\mathbf{K}_i \mathbf{p}) \boldsymbol{\gamma}_i \rangle = (1+\pi) \left\langle -\nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle. \end{aligned} \quad (44)$$

Following (36), for all $i \in [n]$, for all $\mathbf{q} \in \text{cone}_\mu(\mathbf{p}^{mm})$ with $\|\mathbf{q}\| = \|\mathbf{p}^{mm}\|$, $\mathbf{a}' = \mathbf{K} \mathbf{q}$ and $\mathbf{s} = \mathbb{S}(\mathbf{K} \mathbf{p})$, we have found

$$\left| \mathbf{a}'^\top \text{diag}(\mathbf{s}_i) \boldsymbol{\gamma} - \mathbf{a}'^\top \mathbf{s}_i \mathbf{s}_i^\top \boldsymbol{\gamma}_i - \sum_{i \in \mathcal{T}_i} (\mathbf{a}'_{i1} - \mathbf{a}'_{it}) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \right| \leq 2\Gamma A((1 - s_{i1})^2 + Q_i). \quad (45)$$

Plugging in $\mathbf{a}_i, \bar{\mathbf{a}}_i$ in the bound above and assuming $\pi \leq 1$ (w.l.o.g.), (44) is implied by the following stronger inequality

$$\begin{aligned} &-\frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \left(6\Gamma A((1 - s_{i1})^2 + Q_i) + \sum_{i \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \right) \\ &\leq -\frac{1+\pi}{n} \sum_{i=1}^n \ell'_i \cdot \sum_{i \in \mathcal{T}_i} (\bar{\mathbf{a}}_{i1} - \bar{\mathbf{a}}_{it}) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ &= -\frac{1+\pi}{n} \sum_{i=1}^n \ell'_i \cdot \sum_{i \in \mathcal{T}_i} \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}). \end{aligned}$$

First, we claim that $0.5\pi \sum_{i \in \mathcal{T}_i} \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \geq 6\Gamma A((1 - s_{i1})^2 + Q_i)$ for all $i \in [n]$. The proof of this claim directly follows the earlier argument, namely, following (37), (38) and (39) which leads to the choice

$$R \geq \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{C_0 \cdot T\Gamma A}{\pi \gamma_{\min}^{gap}} \right), \quad (46)$$

for some constant $C_0 > 0$. Here, we choose sufficiently large $C_0 \geq 64\pi$ to ensure $R = R_\pi \geq \bar{R}_\mu$.

Following this control over the perturbation term $6\Gamma A((1 - s_{i1})^2 + Q_i)$, to conclude with the result, what remains is proving the comparison

$$-\frac{1}{n} \sum_{i=1}^n \ell'_i \cdot \sum_{t \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \leq -\frac{1 + 0.5\pi}{n} \sum_{i=1}^n \ell'_i \cdot \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}). \quad (47)$$

To proceed, we split the problem into two scenarios.

Scenario 1: $\|\bar{\mathbf{p}} - \mathbf{p}^{mm}\| \leq \epsilon = \frac{\pi}{4A\Theta}$ for some $\epsilon > 0$. In this scenario, for any token, we find that

$$|\mathbf{a}_{it} - \bar{\mathbf{a}}_t| = |\mathbf{k}_{it}^\top (\bar{\mathbf{p}} - \mathbf{p}^{mm})| \leq A\Theta\epsilon = \pi/4.$$

Consequently, we obtain

$$\mathbf{a}_{i1} - \mathbf{a}_{it} \leq \bar{\mathbf{a}}_{i1} - \bar{\mathbf{a}}_{it} + 2A\Theta\epsilon = 1 + 0.5\pi.$$

Similarly, $\mathbf{a}_{i1} - \mathbf{a}_{it} \geq 1 - 0.5\pi \geq 0.5$. Since all terms $\mathbf{a}_{i1} - \mathbf{a}_{it}, s_{it}, \boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}$ in (47) are nonnegative and $(\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \leq (1 + 0.5\pi) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it})$, above implies the desired result (47).

Scenario 2: $\|\bar{\mathbf{p}} - \mathbf{p}^{mm}\| \geq \epsilon = \frac{\pi}{4A\Theta}$. Since $\bar{\mathbf{p}}$ is not (locally) max-margin, in this scenario, for some $i \in [n]$, $\nu = \nu(\epsilon) > 0$, and $\tau \in \mathcal{T}_i$, we have that

$$\bar{\mathbf{p}}^\top (\mathbf{k}_{i1} - \mathbf{k}_{i\tau}) = \mathbf{a}_{i1} - \mathbf{a}_{i\tau} \leq 1 - 2\nu.$$

Here $\tau = \arg \max_{t \in \mathcal{T}_i} \bar{\mathbf{p}}^\top \mathbf{k}_{it}$ denotes the nearest point to $\mathbf{k}_{i1}$ (along the $\bar{\mathbf{p}}$ direction). Note that a non-neighbor $t \in \bar{\mathcal{T}}_i$ cannot be nearest because $\bar{\mathbf{p}} \in \text{cone}_\mu(\mathbf{p}^{mm})$ and (33) holds. Recall that $s_i = \mathbb{S}(\bar{R}\mathbf{a}_i)$ where $\bar{R} = \|\mathbf{p}\|\Theta \geq R\Theta$. To proceed, let $\underline{\mathbf{a}}_i := \min_{t \in \mathcal{T}_i} \mathbf{a}_{i1} - \mathbf{a}_{it}$,

$$\mathcal{I} := \{i \in [n] : \underline{\mathbf{a}}_i \leq 1 - 2\nu\}, \quad [n] - \mathcal{I} := \{i \in [n] : 1 - 2\nu < \underline{\mathbf{a}}_i\}.$$

For all $i \in [n] - \mathcal{I}$,

$$\begin{aligned} \sum_{t \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) - (1 + 0.5\pi) \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ \leq (2A - (1 + 0.5\pi)) \Gamma \sum_{t \in \mathcal{T}_i, \mathbf{a}_{i1} - \mathbf{a}_{it} \geq 1 + \frac{\pi}{2}} s_{it} \\ \leq (2A - (1 + 0.5\pi)) \Gamma T e^{-\bar{R}(1 + \frac{\pi}{2})} \\ \leq 2A\Gamma T e^{-\bar{R}(1 + \frac{\pi}{2})}. \end{aligned} \quad (48)$$

For all $i \in \mathcal{I}$, split the tokens into two groups: Let $\mathcal{N}_i$ be the group of tokens obeying $\mathbf{a}_{i1} - \mathbf{a}_{it} \leq 1 - \nu$ and $\mathcal{T}_i - \mathcal{N}_i$ be the rest of the neighbors. Observe that

$$\frac{\sum_{t \in \mathcal{T}_i - \mathcal{N}_i} s_{it}}{\sum_{t \in \mathcal{T}_i} s_{it}} \leq T \frac{e^{\nu \bar{R}}}{e^{2\nu \bar{R}}} = T e^{-\bar{R}\nu}.$$

Thus, using $|\mathbf{a}_{i1} - \mathbf{a}_{it}| \leq 2A$ and recalling the definition of $\gamma_{\min}^{\text{gap}}$, observe that

$$\sum_{t \in \mathcal{T}_i - \mathcal{N}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \leq \frac{2\Gamma A T e^{-\bar{R}\nu}}{\gamma_{\min}^{\text{gap}}} \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}).$$

Thus,

$$\begin{aligned} \sum_{t \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) &= \sum_{t \in \mathcal{N}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) + \sum_{t \in \mathcal{T}_i - \mathcal{N}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ &\leq \sum_{t \in \mathcal{N}_i} (1 - \nu) s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) + \frac{2\Gamma A T e^{-\bar{R}\nu}}{\gamma_{\min}^{\text{gap}}} \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ &\leq \left(1 - \nu + \frac{2\Gamma A T e^{-\bar{R}\nu}}{\gamma_{\min}^{\text{gap}}}\right) \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ &\leq \left(1 + \frac{2\Gamma A T e^{-\bar{R}\nu}}{\gamma_{\min}^{\text{gap}}}\right) \sum_{t \in \mathcal{T}_i} s_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}). \end{aligned}$$

Hence, choosing

$$R \geq \frac{1}{\nu\Theta} \log \left(\frac{8\Gamma AT}{\gamma_{\min}^{\text{gap}} \pi} \right) \quad (49)$$

results in that

$$\begin{aligned} & \sum_{t \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) - (1 + \frac{\pi}{2}) \sum_{t \in \mathcal{T}_i} \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ & \leq \left(\frac{2\Gamma AT e^{-\bar{R}\nu}}{\gamma_{\min}^{\text{gap}}} - \frac{\pi}{2} \right) \sum_{t \in \mathcal{T}_i} \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ & \leq -\frac{\pi}{4} \sum_{t \in \mathcal{T}_i} \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \\ & \leq -\frac{\pi}{4T} \gamma_{\min}^{\text{gap}} e^{-\bar{R}(1-2\nu)}. \end{aligned} \quad (50)$$

Here, the last inequality follows from the fact that $\sum_{t \in \mathcal{T}_i} \mathbf{s}_{it} \geq \max_{t \in \mathcal{T}_i} \mathbf{s}_{it} \geq \frac{e^{-\bar{R}(1-2\nu)}}{\sum_{t=1}^T e^{-\bar{R}(\mathbf{a}_{i1} - \mathbf{a}_{it})}} \geq e^{-\bar{R}(1-2\nu)}/T$.

From Assumption A, we have $c_{\min} \leq -\ell' \leq c_{\max}$ for some positive constants $c_{\min}$ and $c_{\max}$. It follows from (48) and (50) that

$$\begin{aligned} & -\frac{1}{n} \sum_i \ell'_i \cdot \left(\sum_{t \in \mathcal{T}_i} (\mathbf{a}_{i1} - \mathbf{a}_{it}) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) - \sum_{t \in \mathcal{T}_i} (1 + 0.5\pi) \mathbf{s}_{it} (\boldsymbol{\gamma}_{i1} - \boldsymbol{\gamma}_{it}) \right) \\ & \leq c_{\max} 2\Gamma T \Gamma e^{-\bar{R}(1+\frac{\pi}{2})} - \frac{c_{\min}}{nT} \cdot \frac{\pi \gamma_{\min}^{\text{gap}}}{4} e^{-\bar{R}(1-2\nu)} \\ & \leq 0. \end{aligned}$$

Combing with (49), this is guaranteed by choosing

$$R \geq \max \left\{ \frac{1}{\nu\Theta} \log \left(\frac{8\Gamma AT}{\gamma_{\min}^{\text{gap}} \pi} \right), \frac{1}{(2\nu + \pi/2)\Theta} \log \left(\frac{8n\Gamma AT^2 c_{\max}}{c_{\min} \gamma_{\min}^{\text{gap}} \pi} \right) \right\},$$

where $\nu = \nu(\frac{\pi}{4A\Theta})$ depends only on π and global problem variables.

Combining this with the prior R choice (46) (by taking maximum), we conclude with the statement. $\blacksquare$

B.2 Proof of Theorem 1

Proof. This proof is a direct corollary of Lemma 14 which itself is a special case of the nonlinear head Theorem 8. Let us verify that $f(\mathbf{X}) = \boldsymbol{\nu}^\top \mathbf{X}^\top \mathbb{S}(\mathbf{X}\mathbf{p})$ satisfies the assumptions of Lemma 14 where we replace the nonlinear head with linear $\boldsymbol{\nu}$. To see this, set the optimal sets to be the singletons $\mathcal{O}_i = \{\text{opt}_i\}$. Given $(\mathbf{X}_i, \mathbf{Y}_i)$, let $\mathbf{s}_i = \mathbb{S}(\mathbf{K}_i \mathbf{p})$ and $q_i = q_i^p = \sum_{t \neq \text{opt}_i} \mathbf{s}_{it}$. Recalling score definition $\boldsymbol{\gamma}_i = \mathbf{Y}_i \cdot \mathbf{X}_i \boldsymbol{\nu}$ and setting $\nu_i := \boldsymbol{\gamma}_{i\text{opt}_i}$ and $Z_i := \sum_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it} \mathbf{s}_{it}$, a particular prediction can be written as

$$\begin{aligned} \mathbf{Y}_i \cdot \boldsymbol{\nu}^\top \mathbf{X}_i^\top \mathbb{S}(\mathbf{X}_i \mathbf{p}) &= \boldsymbol{\gamma}_i^\top \mathbf{s}_i = \boldsymbol{\gamma}_{i\text{opt}_i} (1 - q_i) + \sum_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it} \mathbf{s}_{it} \\ &= \nu_i (1 - q_i) + Z_i. \end{aligned}$$

To proceed, we demonstrate the choices for $C, \epsilon > 0$. Let $C := -\min_{i \in [n], t \in [T]} \boldsymbol{\gamma}_{it} \wedge 0$ and $q_{\max} = \max_{i \in [n]} q_i$. Note that $Z_i \geq \sum_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it} \mathbf{s}_{it} \geq q_i \gamma_{\min} \geq -C q_{\max}$. Now, using strict score optimality of opt_i 's for all $i \in [n]$, we set

$$\epsilon := 1 - \sup_{i \in [n]} \frac{\sum_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it} \mathbf{s}_{it}}{\nu_i q_i} \geq 1 - \sup_{i \in [n]} \frac{\sup_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it}}{\boldsymbol{\gamma}_{i\text{opt}_i}} > 0.$$

We conclude by observing $Z_i \leq \nu_i q_i \frac{\sum_{t \neq \text{opt}_i} \boldsymbol{\gamma}_{it} \mathbf{s}_{it}}{\nu_i q_i} \leq \nu_i q_i \epsilon$ as desired. $\blacksquare$

B.3 Proof of Theorem 2

Proof. We first show that $\lim_{t \rightarrow \infty} \|\mathbf{p}(t)\| = \infty$. From Lemma 4, we have

$$\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm\star} \rangle = \frac{1}{n} \sum_{i=1}^n \ell'(Y_i \cdot \mathbf{v}^\top X_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p})) \cdot \langle \mathbf{K}_i \mathbf{p}^{mm\star}, \mathbb{S}'(\mathbf{a}_i) \boldsymbol{\gamma}_i \rangle,$$

where $\boldsymbol{\gamma}_i = Y_i \cdot X_i \mathbf{v}$ and $\mathbf{a}_i = \mathbf{K}_i \mathbf{p}$.

It follows from Lemma 4 that $\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm\star} \rangle < 0$ for all $\mathbf{p} \in \mathbb{R}^d$. Hence, for any finite $\mathbf{p}$, $\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm\star} \rangle$ cannot be equal to zero, as a sum of negative terms. Therefore, there are no finite critical points $\mathbf{p}$, for which $\nabla \mathcal{L}(\mathbf{p}) = 0$. On the other hand, Lemma 6 states $\nabla \mathcal{L}(\mathbf{p}(t)) \rightarrow 0$ which implies that $\|\mathbf{p}(t)\| \rightarrow \infty$.

Next, we provide the directional convergence for the setting $n = 1$. Let us consider an arbitrary value of $\epsilon \in (0, 1)$ and set $\pi = \epsilon/(1 - \epsilon)$. As $\lim_{t \rightarrow \infty} \|\mathbf{p}(t)\| = \infty$, we can select a specific t_ϵ such that for all $t \geq t_\epsilon$, it holds that $\|\mathbf{p}(t)\| \geq R_\epsilon \vee 1/2$ for any choice of R_ϵ . To proceed, we choose R_ϵ based on Lemma 5 so that for any $t \geq t_\epsilon$, we have that

$$\left\langle -\nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}^{mm\star}}{\|\mathbf{p}^{mm\star}\|} \right\rangle \geq (1 - \epsilon) \left\langle -\nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|} \right\rangle.$$

Multiplying both sides by the stepsize η and using the gradient descent update, we get

$$\begin{aligned} \left\langle \mathbf{p}(t+1) - \mathbf{p}(t), \frac{\mathbf{p}^{mm\star}}{\|\mathbf{p}^{mm\star}\|} \right\rangle &\geq (1 - \epsilon) \left\langle \mathbf{p}(t+1) - \mathbf{p}(t), \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|} \right\rangle \\ &= \frac{(1 - \epsilon)}{2\|\mathbf{p}(t)\|} \left(\|\mathbf{p}(t+1)\|^2 - \|\mathbf{p}(t)\|^2 - \|\mathbf{p}(t+1) - \mathbf{p}(t)\|^2 \right) \\ &\geq (1 - \epsilon) \left(\frac{1}{2\|\mathbf{p}(t)\|} \left(\|\mathbf{p}(t+1)\|^2 - \|\mathbf{p}(t)\|^2 \right) - \|\mathbf{p}(t+1) - \mathbf{p}(t)\|^2 \right) \quad (51) \\ &\geq (1 - \epsilon) \left(\|\mathbf{p}(t+1)\| - \|\mathbf{p}(t)\| - \|\mathbf{p}(t+1) - \mathbf{p}(t)\|^2 \right) \\ &\geq (1 - \epsilon) \left(\|\mathbf{p}(t+1)\| - \|\mathbf{p}(t)\| - 2\eta (\mathcal{L}(\mathbf{p}(t)) - \mathcal{L}(\mathbf{p}(t+1))) \right). \end{aligned}$$

Here, the second inequality is obtained from $\|\mathbf{p}(t)\| \geq 1/2$; the third inequality follows since for any $a, b > 0$, we have $(a^2 - b^2)/(2b) - (a - b) \geq 0$; and the last inequality uses Lemma 6.

Summing the above inequality over $t \geq t_\epsilon$ gives

$$\left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|}, \frac{\mathbf{p}^{mm\star}}{\|\mathbf{p}^{mm\star}\|} \right\rangle \geq 1 - \epsilon + \frac{C(\epsilon, \eta)}{\|\mathbf{p}(t)\|},$$

for some finite constant $C(\epsilon, \eta)$ defined as

$$C(\epsilon, \eta) := \left\langle \mathbf{p}(t_\epsilon), \frac{\mathbf{p}^{mm\star}}{\|\mathbf{p}^{mm\star}\|} \right\rangle - (1 - \epsilon)\|\mathbf{p}(t_\epsilon)\| - 2\eta(1 - \epsilon)(\mathcal{L}(\mathbf{p}(t_\epsilon)) - \mathcal{L}^*), \quad (52)$$

where $\mathcal{L}^*$ denotes the minimum objective.

Since $\|\mathbf{p}(t)\| \rightarrow \infty$, we get

$$\liminf_{t \rightarrow \infty} \left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|}, \frac{\mathbf{p}^{mm\star}}{\|\mathbf{p}^{mm\star}\|} \right\rangle \geq 1 - \epsilon.$$

Given that we can choose any value of $\epsilon \in (0, 1)$, we have $\mathbf{p}(t)/\|\mathbf{p}(t)\| \rightarrow \mathbf{p}^{mm\star}/\|\mathbf{p}^{mm\star}\|$. ■

B.4 Proof of Theorem 3

Proof. Following the proof of Lemma 7, let $(\mathcal{T}_i)_{i=1}^n$ denote the sets of SVM-neighbors as defined in Definition 2. We define $\bar{\mathcal{T}}_i = [T] - \mathcal{T}_i - \{\alpha_i\}$ as the tokens that are non-SVM neighbors. Additionally, let μ be defined as in (32). Let us denote the initialization lower bound as $R_\mu^0 := R$, where R is given in the Theorem 3's statement. Consider an arbitrary value of $\epsilon \in (0, \mu/2)$ and let $1/(1 + \pi) = 1 - \epsilon$.

We additionally denote $R_\epsilon \leftarrow R_\pi \vee 1/2$ where R_π was defined in Lemma 7(L3.). At initialization $\mathbf{p}(0)$, we set $\epsilon = \mu/2$ to obtain $R_\mu^0 = R_{\mu/2}$ and provide the proof in four steps:

Step 1: There are no stationary points within $C_{\mu, R_\mu^0}(\mathbf{p}^{mm})$. We begin by proving that there are no stationary points within $C_{\mu, R_\mu^0}(\mathbf{p}^{mm})$. Then, since $R_\mu^0 \geq \bar{R}_\mu$ per Lemma 7, we can apply (L2.) to find that: For all $\mathbf{q}, \mathbf{p} \in \text{cone}_\mu(\mathbf{p}^{mm})$ with $\mathbf{q} \neq 0$ and $\|\mathbf{p}\| \geq R_\mu^0$, we have that $-\mathbf{q}^\top \nabla \mathcal{L}(\mathbf{p})$ is strictly positive.

Step 2: It follows from Lemma 7(L3.) that, for all $\epsilon \in (0, \mu/2)$, all $\mathbf{p} \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$ satisfy

$$\left\langle -\nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \geq (1 - \epsilon) \left\langle -\nabla \mathcal{L}(\mathbf{p}), \frac{\mathbf{p}}{\|\mathbf{p}\|} \right\rangle. \quad (53)$$

The argument above applies to a general $\epsilon \in (0, \mu/2)$. However, at initialization $\mathbf{p}(0)$, we set $\epsilon = \mu/2$ to obtain our earlier R_μ^0 choice. To proceed, for any $\epsilon \in (0, \mu/2)$, we will show that after gradient descent enters the conic set $C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$ for the first time, it will never leave the set. Let t_ϵ be the first time gradient descent enters $C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$. In Step 4, we will prove that such t_ϵ is guaranteed to exist. Additionally, for $\epsilon \leftarrow \mu/2$, note that $t_\epsilon = 0$ i.e. the point of initialization.

Step 3: Updates remain inside the cone $C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$. By leveraging the results from Step 1 and Step 2, we demonstrate that the gradient iterates, with an appropriate constant step size, starting from $\mathbf{p}(t_\epsilon) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$, remain within this cone.

We proceed by induction. Suppose that the claim holds up to iteration $t \geq t_\epsilon$. This implies that $\mathbf{p}(t) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$. Hence, recalling cone definition, for $\mu > 0$ and R_ϵ , we have $\left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \geq 1 - \mu$ and $\|\mathbf{p}(t)\| \geq R_\epsilon$. Let

$$\rho(t) := -\frac{1}{1 - \epsilon} \left\langle \nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle.$$

Note that $\rho(t) > 0$ due to Step 1. This together with the gradient descent update rule gives

$$\begin{aligned} \left\langle \frac{\mathbf{p}(t+1)}{\|\mathbf{p}(t+1)\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle &= \left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|} - \frac{\eta}{\|\mathbf{p}(t)\|} \nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \\ &\geq 1 - \mu - \frac{\eta}{\|\mathbf{p}(t)\|} \left\langle \nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \\ &= 1 - \mu + \frac{\eta \rho(t)(1 - \epsilon)}{\|\mathbf{p}(t)\|}. \end{aligned} \quad (54a)$$

Note that from Lemma 7, we have $\langle \nabla \mathcal{L}(\mathbf{p}(t)), \mathbf{p}(t) \rangle < 0$ which implies that $\|\mathbf{p}(t+1)\| \geq \|\mathbf{p}(t)\|$. This together with R_ϵ definition and $\|\mathbf{p}(t)\| \geq 1/2$ implies that

$$\begin{aligned} \|\mathbf{p}(t+1)\| &\leq \frac{1}{2\|\mathbf{p}(t)\|} (\|\mathbf{p}(t+1)\|^2 + \|\mathbf{p}(t)\|^2) \\ &= \frac{1}{2\|\mathbf{p}(t)\|} (2\|\mathbf{p}(t)\|^2 - 2\eta \langle \nabla \mathcal{L}(\mathbf{p}(t)), \mathbf{p}(t) \rangle + \eta^2 \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2) \\ &\leq \|\mathbf{p}(t)\| - \frac{\eta}{\|\mathbf{p}(t)\|} \langle \nabla \mathcal{L}(\mathbf{p}(t)), \mathbf{p}(t) \rangle + \eta^2 \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2. \end{aligned} \quad (54b)$$

Hence, using (53)

$$\begin{aligned} \frac{\|\mathbf{p}(t+1)\|}{\|\mathbf{p}(t)\|} &\leq 1 - \frac{\eta}{\|\mathbf{p}(t)\|} \left\langle \nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|} \right\rangle + \eta^2 \frac{\|\nabla \mathcal{L}(\mathbf{p}(t))\|^2}{\|\mathbf{p}(t)\|} \\ &\leq 1 - \frac{\eta}{(1 - \epsilon)\|\mathbf{p}(t)\|} \left\langle \nabla \mathcal{L}(\mathbf{p}(t)), \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle + \eta^2 \frac{\|\nabla \mathcal{L}(\mathbf{p}(t))\|^2}{\|\mathbf{p}(t)\|} \\ &= 1 + \frac{\eta \rho(t)}{\|\mathbf{p}(t)\|} + \frac{\eta^2 \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2}{\|\mathbf{p}(t)\|} =: C_1(\rho(t), \eta). \end{aligned} \quad (54c)$$

Here, the second inequality follows from (53).

Now, it follows from (54a) and (54c) that

$$\begin{aligned}
\left\langle \frac{\mathbf{p}(t+1)}{\|\mathbf{p}(t+1)\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle &\geq \frac{1}{C_1(\rho(t), \eta)} \left(1 - \mu + \frac{\eta \rho(t)(1-\epsilon)}{\|\mathbf{p}(t)\|} \right) \\
&= 1 - \mu + \frac{1}{C_1(\rho(t), \eta)} \left((1-\mu)(1 - C_1(\rho(t), \eta)) + \frac{\eta \rho(t)(1-\epsilon)}{\|\mathbf{p}(t)\|} \right) \\
&= 1 - \mu + \frac{\eta}{C_1(\rho(t), \eta)} \left((\mu-1) \left(\frac{\rho(t)}{\|\mathbf{p}(t)\|} + \frac{\eta \|\nabla \mathcal{L}(\mathbf{p}(t))\|^2}{\|\mathbf{p}(t)\|} \right) + \frac{\rho(t)(1-\epsilon)}{\|\mathbf{p}(t)\|} \right) \quad (55) \\
&= 1 - \mu + \frac{\eta}{C_1(\rho(t), \eta)} \left(\frac{\rho(t)(\mu-\epsilon)}{\|\mathbf{p}(t)\|} - \eta(1-\mu) \frac{\|\nabla \mathcal{L}(\mathbf{p}(t))\|^2}{\|\mathbf{p}(t)\|} \right) \\
&\geq 1 - \mu,
\end{aligned}$$

where the last inequality uses our choice of stepsize $\eta \leq 1/L_p$ in Theorem 3's statement. Specifically, we need η to be small to ensure the last inequality. We will guarantee this by choosing a proper R_ϵ in Lemma 7. Specifically, Lemma 7 leaves the choice of C_0 in R_ϵ lower bound of (46) open (it can always be chosen larger). Here, by choosing $C_0 \gtrsim 1/L_p$ will ensure $\eta \leq 1/L_p$ works well.

To proceed, we have that

$$\begin{aligned}
\frac{(\mu-\epsilon)}{1-\mu} \frac{\rho(t)}{\|\nabla \mathcal{L}(\mathbf{p}(t))\|^2} &\geq \frac{\mu-\epsilon}{1-\mu} \frac{1}{1-\epsilon} \frac{c}{C} \frac{\Theta}{\bar{A}} \frac{1}{\bar{A}CT} e^{R_\mu^0 \Theta/2} \\
&\geq \frac{\mu}{2(1-\mu)(1-\frac{\mu}{2})} \frac{c}{C} \frac{\Theta}{\bar{A}} \frac{1}{\bar{A}CT} e^{R_\mu^0 \Theta/2} \geq \eta. \quad (56)
\end{aligned}$$

Here, the second inequality uses our choice of $\epsilon \in (0, \mu/2)$ (see **Step 2**), and the first inequality is obtained from Lemma 7 since

$$\begin{aligned}
\frac{\rho(t)}{\|\nabla \mathcal{L}(\mathbf{p}(t))\|} &= -\frac{1}{1-\epsilon} \left\langle \frac{\nabla \mathcal{L}(\mathbf{p}(t))}{\|\nabla \mathcal{L}(\mathbf{p}(t))\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \geq \frac{1}{1-\epsilon} \frac{c}{C} \frac{\Theta}{\bar{A}}, \\
\frac{1}{\|\nabla \mathcal{L}(\mathbf{p}(t))\|} &\geq \frac{1}{\bar{A}C \frac{1}{n} \sum_{i=1}^n (1-s_{i\alpha_i})} \geq \frac{1}{\bar{A}CT e^{-R_\mu^0 \Theta/2}}
\end{aligned}$$

for some data dependent constants c and C , $\bar{A} = \max_{i \in [n], t, \tau \in [T]} \|\mathbf{k}_{it} - \mathbf{k}_{i\tau}\|$, and $\Theta = 1/\|\mathbf{p}^{mm}\|$.

Next, we will demonstrate that the choice of η in (56) does indeed meet our step size condition as stated in the theorem, i.e., $\eta \leq 1/L_p$. Recall that $1/(1+\pi) = 1-\epsilon$, which implies that $\pi = \epsilon/(1-\epsilon)$. Combining this with (46), we obtain:

$$\begin{aligned}
R_\pi &\geq \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{C_0 T \Gamma A}{\pi \gamma_{\min}^{gap}} \right), \quad \text{where } C_0 \geq 64\pi, \\
\Rightarrow R_\epsilon &\geq \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{(1-\epsilon)C_0 T \Gamma A}{\epsilon \gamma_{\min}^{gap}} \right), \quad \text{where } C_0 \geq 64 \frac{\epsilon}{1-\epsilon}.
\end{aligned}$$

On the other hand, at the initialization, we have $\epsilon = \mu/2$ which implies that

$$R_\mu^0 \geq \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{(2-\mu)C_0 T \Gamma A}{\mu \gamma_{\min}^{gap}} \right), \quad \text{where } C_0 \geq 64 \frac{\mu}{(2-\mu)}. \quad (57)$$

In the following, we will determine a lower bound on C_0 such that our step size condition in Theorem 3's statement, i.e., $\eta \leq 1/L_p$, is satisfied. Note that for the choice of η in (56) to meet the condition $\eta \leq 1/L_p$, the following condition must hold:

$$\frac{1}{L_p} \leq \frac{\mu}{(2-\mu)} \frac{1}{C_2 T} e^{R_\mu^0 \Theta/2} \Rightarrow R_\mu^0 \geq \frac{2}{\Theta} \log \left(\frac{1}{L_p} \frac{(2-\mu)}{\mu} C_2 T \right), \quad (58)$$

where $C_2 = (1-\mu) \frac{\bar{A}^2 C^2}{\Theta c}$.

This together with (57) implies that for sufficiently large

$$R_\mu^0 \geq \frac{\max(2, \delta^{-1})}{\Theta} \log \left(\frac{(2-\mu)C_3 T}{\mu} \right), \quad \text{where } C_3 = \frac{C_0 \Gamma A}{\gamma_{\min}^{gap}} \vee \frac{C_2}{L_p},$$

the step size bound in (56) ensures that $\eta \leq 1/L_p$ guarantees (55). Hence, $\mathbf{p}(t+1)$ remains within the cone, i.e., $\mathbf{p}(t+1) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$.

Step 4: The correlation of $\mathbf{p}(t)$ and $\mathbf{p}^{mm}$ increases over t . The remainder is similar to the proof of Theorem 2. From Step 3, we have that all iterates remain within the initial conic set i.e. $\mathbf{p}(t) \in C_{\mu, R_\mu^0}(\mathbf{p}^{mm})$ for all $t \geq 0$. Note that it follows from Lemma 7 that $\langle \nabla \mathcal{L}(\mathbf{p}), \mathbf{p}^{mm}/\|\mathbf{p}^{mm}\| \rangle < 0$, for any finite $\mathbf{p} \in C_{\mu, R_\mu^0}(\mathbf{p}^{mm})$. Hence, there are no finite critical points $\mathbf{p} \in C_{\mu, R_\mu^0}(\mathbf{p}^{mm})$, for which $\nabla \mathcal{L}(\mathbf{p}) = 0$. Now, based on Lemma 6, which guarantees that $\nabla \mathcal{L}(\mathbf{p}(t)) \rightarrow 0$, this implies that $\|\mathbf{p}(t)\| \rightarrow \infty$. Consequently, for any choice of $\epsilon \in (0, \mu/2)$ there is a time t_ϵ such that, for all $t \geq t_\epsilon$, $\mathbf{p}(t) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$. Once within $C_{\mu, R_\epsilon}(\mathbf{p}^{mm})$, following similar steps in (51) and (52), for any $t \geq t_\epsilon$,

$$\left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \geq 1 - \epsilon + \frac{C_2(\epsilon, \eta)}{\|\mathbf{p}(t)\|}, \quad \mathbf{p}(t) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm}),$$

for some finite constant $C_2(\epsilon, \eta)$. Consequently,

$$\liminf_{t \rightarrow \infty} \left\langle \frac{\mathbf{p}(t)}{\|\mathbf{p}(t)\|}, \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|} \right\rangle \geq 1 - \epsilon, \quad \text{where } \mathbf{p}(t) \in C_{\mu, R_\epsilon}(\mathbf{p}^{mm}).$$

Since the choice of $\epsilon \in (0, \mu/2)$ is arbitrary, we obtain $\mathbf{p}(t)/\|\mathbf{p}(t)\| \rightarrow \mathbf{p}^{mm}/\|\mathbf{p}^{mm}\|$. $\blacksquare$

B.5 Proof of Theorem 4

B.5.1 Supporting Lemma

We present a lemma that will aid in simplifying our analysis. We begin with a definition.

Definition 3 (Selected-tokens, Neighbors, Margins, and Neighbor-optimality of a direction)

Let $\mathbf{q} \in \mathbb{R}^d - \{\mathbf{0}\}$ and $(Y_i, \mathbf{K}_i, X_i)_{i=1}^n$ be our dataset. We define the (possibly non-unique) selected-tokens of $\mathbf{q}$ as follows:³

$$\alpha_i \in \arg \max_{t \in [T]} \mathbf{k}_{it}^\top \mathbf{q}. \quad (59)$$

Next, we define the margin and directional-neighbors for $\mathbf{q}$ as the minimum margin tokens to the selected-tokens, i.e.,

$$\Gamma_{\mathbf{q}} = \min_{i \in [n], t \neq \alpha_i} (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{q}, \quad (60)$$

$$\mathcal{M}_{\mathbf{q}} = \{(i, t) \mid (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{q} = \Gamma_{\mathbf{q}}\}. \quad (61)$$

Finally, we say that $\mathbf{q}$ is neighbor-optimal if the scores of its directional-neighbors are strictly less than the corresponding selected-token. Concretely, for all $(i, t) \in \mathcal{M}_{\mathbf{q}}$, we require that

$$\gamma_{it} = Y_i \cdot \mathbf{x}_{it}^\top \mathbf{v} < \gamma_{i\alpha_i} = Y_i \cdot \mathbf{x}_{i\alpha_i}^\top \mathbf{v}.$$

Lemma 8 (When does one direction dominate another?) Suppose $\mathbf{q}, \mathbf{p} \in \mathbb{R}^d$ be two unit Euclidean norm vectors with identical selected tokens. Specifically, for each $i \in [n]$, there exists unique $\alpha_i \in [T]$ such that $\alpha_i = \arg \max_{t \in [T]} \mathbf{k}_{it}^\top \mathbf{q} = \arg \max_{t \in [T]} \mathbf{k}_{it}^\top \mathbf{p}$. Suppose directional margins obey $\Gamma_{\mathbf{q}} < \Gamma_{\mathbf{p}}$ and set $\delta_\Gamma = \Gamma_{\mathbf{p}} - \Gamma_{\mathbf{q}}$.

- Suppose $\mathbf{q}$ and $\mathbf{p}$ are both neighbor-optimal. Then, for some $R(\delta_\Gamma)$ and all $R > R(\delta_\Gamma)$, we have that $\mathcal{L}(R \cdot \mathbf{p}) < \mathcal{L}(R \cdot \mathbf{q})$.
- Suppose $\mathbf{q}$ has a unique directional-neighbor and is not neighbor-optimal (i.e. this neighbor has higher score). Let δ_q be the margin difference between unique directional-neighbor and the second-most minimum-margin neighbor (i.e. the one after the unique one, see (65)) of $\mathbf{q}$. Then, for some $R(\delta_\Gamma \wedge \delta_q)$ and all $R > R(\delta_\Gamma \wedge \delta_q)$, we have that $\mathcal{L}(R \cdot \mathbf{q}) < \mathcal{L}(R \cdot \mathbf{p})$.

Proof. We prove these two statements in order. First define the directional risk baseline induced by letting $R \rightarrow \infty$ and purely selecting the tokens $\alpha = (\alpha_i)_{i=1}^n$. This is given by

$$\mathcal{L}_\star := \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot \mathbf{v}^\top \mathbf{x}_{i\alpha_i}).$$

³If α_i is unique for all $i \in [n]$, let us call it, unique selected tokens.

We evaluate $\mathbf{q}, \mathbf{p}$ with respect to $\mathcal{L}_\star$. To proceed, let $s_i = \mathbb{S}(R\mathbf{K}_i\mathbf{q})$. Define $\Gamma_q^{it} = \mathbf{k}_{i\alpha_i}^\top \mathbf{q} - \mathbf{k}_{it}^\top \mathbf{q}$. Note that, the smallest value for $t \neq \alpha_i$ is achieved for Γ_q . For sufficiently large $R \gtrsim O(\log(T)/\Gamma_q)$, observe that, for $t \neq \alpha_i$

$$e^{-R\Gamma_q^{it}} \geq s_{it} = \frac{e^{R\mathbf{k}_{it}^\top \mathbf{q}}}{\sum_{t \in [T]} e^{R\mathbf{k}_{it}^\top \mathbf{q}}} \geq 0.5e^{-R\Gamma_q^{it}}. \quad (62)$$

Recalling the score definition and let M_+, M_- be the upper and lower bounds on $-\ell'$ over its bounded domain that scores fall on, respectively. Note that, for some intermediate $M_+ \geq M_i \geq M_-$ values, we have

$$\begin{aligned} \mathcal{L}(R\mathbf{q}) - \mathcal{L}_\star &= \frac{1}{n} \sum_{i=1}^n \ell\left(\sum_{t \in [T]} s_{it} \gamma_{it}\right) - \ell(\gamma_{i\alpha_i}) \\ &= \frac{1}{n} \sum_{i=1}^n M_i \sum_{t \neq \alpha_i} s_{it} (\gamma_{i\alpha_i} - \gamma_{it}). \end{aligned}$$

Now, using (62) for a refreshed $M_+ \geq M_{it} \geq 0.5M_-$ values, we can write

$$\mathcal{L}(R\mathbf{q}) - \mathcal{L}_\star = \frac{1}{n} \sum_{i=1}^n \sum_{t \neq \alpha_i} M_{it} e^{-R\Gamma_q^{it}} (\gamma_{i\alpha_i} - \gamma_{it}). \quad (63)$$

The same bound also applies to $\mathbf{p}$ with some M'_{it} , multipliers

$$\mathcal{L}(R\mathbf{p}) - \mathcal{L}_\star = \frac{1}{n} \sum_{i=1}^n \sum_{t \neq \alpha_i} M'_{it} e^{-R\Gamma_p^{it}} (\gamma_{i\alpha_i} - \gamma_{it}). \quad (64)$$

We can now proceed with the proof.

Case 1: $\mathbf{q}$ and $\mathbf{p}$ are both neighbor-optimal. This means that $\gamma_{i\alpha_i} - \gamma_{it} > 0$ for all $i \in [n], t \neq \alpha_i$. Let $K_+ > K_- > 0$ be upper and lower bounds on $\gamma_{i\alpha_i} - \gamma_{it}$ values. We can now upper bound the right hand side of (63) via

$$M_+ K_+ T e^{-R\Gamma_q} \geq \mathcal{L}(R\mathbf{q}) - \mathcal{L}_\star \geq \frac{1}{2n} M_- K_- e^{-R\Gamma_q}.$$

Consequently, $\mathcal{L}(R\mathbf{q}) > \mathcal{L}(R\mathbf{p})$ as soon as $\frac{1}{2n} M_- K_- e^{-R\Gamma_q} > M_+ K_+ T e^{-R\Gamma_p}$. Since M_+, K_+, n, T are global constants, this happens under the stated condition on the margin gap $\Gamma_p - \Gamma_q$.

Case 2: $\mathbf{q}$ has a unique directional-neighbor and is not neighbor-optimal. In this scenario, $\mathcal{L}(R\mathbf{q}) - \mathcal{L}_\star$ is actually negative for large R . To proceed, define the maximum score difference $K_+ = \sup_{i, t \neq \alpha_i} |\gamma_{i\alpha_i} - \gamma_{it}|$. Also let (j, β) be the unique directional neighbor achieving the minimum margin Γ_q . Then, δ_q – the margin difference between unique directional-neighbor and the second minimum-margin neighbor (i.e. the one after the unique one) of $\mathbf{q}$ – is defined as

$$\delta_q = \min_{i \in [n], t \neq \alpha_i, (i, t) \neq (j, \beta)} \Gamma_q^{it} - \Gamma_q. \quad (65)$$

To proceed, we can write

$$\mathcal{L}(R\mathbf{p}) - \mathcal{L}_\star \geq -M_+ K_+ T e^{-R\Gamma_p}.$$

On the other hand, setting $\kappa = \gamma_{j\beta} - \gamma_{j\alpha_j} > 0$, we can bound

$$\begin{aligned} \mathcal{L}(R\mathbf{q}) - \mathcal{L}_\star &= -\frac{1}{n} M_{j\beta} e^{-R\Gamma_q} \kappa + \frac{1}{n} \sum_{i \in [n], t \neq \alpha_i, (i, t) \neq (j, \beta)} M_{it} e^{-R\Gamma_q^{it}} (\gamma_{i\alpha_i} - \gamma_{it}) \\ &\leq -\frac{1}{n} M_- e^{-R\Gamma_q} \kappa + M_+ K_+ T e^{-R(\Gamma_q + \delta_q)}. \end{aligned}$$

Consequently, we have found that $\mathcal{L}(R\mathbf{p}) > \mathcal{L}(R\mathbf{q})$ as soon as

$$\frac{1}{n} M_- e^{-R\Gamma_q} \kappa \geq M_+ K_+ T (e^{-R(\Gamma_q + \delta_q)} + e^{-R\Gamma_p})$$

This happens when $R \gtrsim \frac{1}{\delta_q \wedge (\Gamma_p - \Gamma_q)}$ (up to logarithmic terms) establishing the desired statement. $\blacksquare$

B.5.2 Proof of Theorem 4

Define the locally-optimal unit directions

$$\mathcal{P}^{mm} = \left\{ \frac{\mathbf{p}^{mm}(\alpha)}{\|\mathbf{p}^{mm}(\alpha)\|} \mid \alpha \text{ is a locally-optimal set of indices} \right\}.$$

The theorem below shows that cone-restricted regularization paths can only directionally converge to an element of this set.

Theorem 7 (Non-LOMM Regularization Paths Fail) *Fix a unit Euclidean norm vector $\mathbf{q} \in \mathbb{R}^d$ such that $\mathbf{q} \notin \mathcal{P}^{mm}$. Assume that the token scores are distinct (i.e., $\gamma_{it} \neq \gamma_{i\tau}$ for $t \neq \tau$) and the key embeddings $\mathbf{k}_{it}$ are in general position. Specifically, we require the following conditions to hold⁴:*

- When $m = d$, all matrices $\tilde{\mathbf{K}} \in \mathbb{R}^{m \times d}$ where each row of $\tilde{\mathbf{K}}$ has the form $\mathbf{k}_{it} - \mathbf{k}_{i\alpha_i}$ for a unique $(i, \alpha_i, t \neq \alpha_i)$ tuple, are full-rank.
- When $m = d + 1$, the vector of all ones is not in the range space of any such $\tilde{\mathbf{K}}$ matrix.

Fix arbitrary $\epsilon > 0, R_0 > 0$. Define the local regularization path of $\mathbf{q}$ as its (ϵ, R_0) -conic neighborhood:

$$\bar{\mathbf{p}}(R) = \arg \min_{\mathbf{p} \in C_{\epsilon, R_0}(\mathbf{q}), \|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p}), \quad \text{where } C_{\epsilon, R_0}(\mathbf{q}) = \text{cone}_\epsilon(\mathbf{q}) \cap \{\mathbf{p} \in \mathbb{R}^d \mid \|\mathbf{p}\| \geq R_0\}. \quad (66)$$

Then, either $\lim_{R \rightarrow \infty} \|\bar{\mathbf{p}}(R)\| < \infty$ or $\lim_{R \rightarrow \infty} \bar{\mathbf{p}}(R)/\|\bar{\mathbf{p}}(R)\| \neq \mathbf{q}$. In both scenarios $\lim_{R \rightarrow \infty} \bar{\mathbf{p}}(R)/R \neq \mathbf{q}$.

Proof. We will prove the result by dividing the problem into distinct cases. In each case, we will construct an alternative direction that achieves a strictly better objective than some $\delta = \delta(\epsilon) > 0$ neighborhood of $\mathbf{q}$, thereby demonstrating the suboptimality of the $\mathbf{q}$ direction. Let's define the δ neighborhood as follows:

$$\mathcal{N}_\delta = \left\{ \mathbf{p} \mid \left\| \frac{\mathbf{p}}{\|\mathbf{p}\|} - \mathbf{q} \right\| \leq \delta \quad \text{and} \quad \|\mathbf{p}\| \geq R_0 \right\}. \quad (67)$$

Now, let's recall a few more definitions based on Definition 3. First, the tokens selected by $\mathbf{q}$ are given by (59). To proceed, let's initially consider the scenario where α_i is unique for all $i \in [n]$, meaning that as we let $c \rightarrow \infty$, $c \cdot \mathbf{q}$ will choose a single token per input. Later, we will revisit the setting when $\arg \max$ is not a singleton, and $\mathbf{q}$ is allowed to select multiple tokens.

Additionally, it's important to note that $\|\bar{\mathbf{p}}(R)\|$ is non-decreasing by definition. Suppose it has a finite upper bound $\|\bar{\mathbf{p}}(R)\| \leq M$ for all $R < \infty$. In that scenario, we have $\lim_{R \rightarrow \infty} \frac{\bar{\mathbf{p}}(R)}{R} = 0 \neq \mathbf{q}$.

• **(A) $\mathbf{q}$ selects a single token per input:** Given that the indices $\alpha = (\alpha_i)_{i=1}^n$ defined in (59) are uniquely determined, we can conclude that the $\mathbf{q}$ direction eventually selects tokens $\mathbf{k}_{i\alpha_i}$. Recall the definition of the margin $\Gamma_{\mathbf{q}}$ from (60) and the set of directional neighbors, which is defined as the indices that achieve $\Gamma_{\mathbf{q}}$, as shown in (61). Let us refer to $\mathbf{q}$ as *neighbor-optimal* if $\gamma_{it} < \gamma_{i\alpha_i}$ for all $(i, t) \in \mathcal{M}_{\mathbf{q}}$.

We will consider two cases for this scenario: when $\mathbf{q}$ is neighbor-optimal and when $\mathbf{q}$ is not neighbor-optimal.

◊ **(A1) $\mathbf{q}$ is neighbor-optimal.** In this case, we will argue that max-margin direction $\bar{\mathbf{p}}^{mm} := \mathbf{p}^{mm}(\alpha)/\|\mathbf{p}^{mm}(\alpha)\|$ can be used to construct a strictly better objective than $\mathbf{q}$. Note that $\mathbf{p}^{mm}(\alpha)$ exists because $\mathbf{q}$ is already a viable separating direction for tokens α . Specifically, consider the direction $\mathbf{q}' = \frac{\mathbf{q} + \epsilon \bar{\mathbf{p}}^{mm}}{\|\mathbf{q} + \epsilon \bar{\mathbf{p}}^{mm}\|}$. Observe that, $\mathbf{q}'$ lies within $\text{cone}_{2\epsilon}(\mathbf{q})$,⁵

$$\mathbf{q}^\top \mathbf{q}' \geq \frac{1 - \epsilon}{1 + \epsilon} \geq 1 - 2\epsilon.$$

We now argue that, there exists $\delta = \delta_\epsilon > 0$ such that for all $R > R_\epsilon$

$$\min_{R_\epsilon \leq r \leq R} \mathcal{L}(r \cdot \mathbf{q}') < \min_{\mathbf{p} \in \mathcal{N}_\delta, R_\epsilon \leq \|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p}). \quad (68)$$

⁴This requirement holds for general data because it is guaranteed by adding arbitrarily small independent gaussian perturbations to keys $\mathbf{k}_{it}$.

⁵As a result, let us prove the result for $\epsilon \leftarrow 2\epsilon$ without losing generality.

To prove this, we study the margin $\Gamma_{q'}$ induced by q' and the maximum margin Γ_δ induced within $p \in \mathcal{N}_\delta$. Concretely, we will show that $\Gamma_{q'} > \Gamma_\delta$ and directly apply the first statement of Lemma 8 to conclude with (68).

Let $\Gamma = 1/\|p^{mm}(\alpha)\|$ be the margin induced by p^{mm} . Note that $\Gamma > \Gamma_q$ by the optimality of $p^{mm}(\alpha)$ and the fact that $q \neq p^{mm}(\alpha)$. Consequently, we can lower and upper bound the margins via

$$\begin{aligned}\Gamma_{q'} &= \min_{i \in [n]} \min_{t \neq \alpha_i} (k_{i\alpha_i} - k_{it})^\top q' \geq \frac{\Gamma_q + \epsilon \Gamma}{1 + \epsilon} \geq \Gamma_q + \frac{\epsilon}{2}(\Gamma - \Gamma_q), \\ \Gamma_\delta &= \max_{p \in \mathcal{N}_\delta} \min_{i \in [n]} \min_{t \neq \alpha_i} (k_{i\alpha_i} - k_{it})^\top p / \|p\| \\ &\leq \max_{\|p\| \leq 1} \min_{i \in [n]} \min_{t \neq \alpha_i} (k_{i\alpha_i} - k_{it})^\top (q + \delta r) \\ &\leq \Gamma_q + M\delta,\end{aligned}$$

where $M = \max_{i,t,\tau} \|k_{it} - k_{i\tau}\|$.

Consequently, setting $\delta = \frac{\epsilon}{4M}(\Gamma - \Gamma_q)$, we find that

$$\Gamma_{q'} \geq \Gamma_\delta + \frac{\epsilon}{4}(\Gamma - \Gamma_q).$$

Equipped with this inequality, we apply the first statement of Lemma 8 which concludes that⁶ for some $R_\epsilon = R(\frac{\epsilon}{4}(\Gamma - \Gamma_q))$ and all $R > R_\epsilon$, (68) holds. This in turn implies that, within C_ϵ , the optimal solution is

- either upper bounded by R_ϵ in ℓ_2 norm (i.e. $\lim_{R \rightarrow \infty} \|\bar{p}(R)\| < \infty$) or
- at least $\delta = \delta(\epsilon) > 0$ away from q after ℓ_2 -normalization i.e. $\|\frac{\bar{p}(R)}{\|\bar{p}(R)\|} - q\| \geq \delta$.

In either scenario, we have proven that $\lim_{R \rightarrow \infty} \frac{\bar{p}(R)}{R} \neq q$.

◊ **(A2) q is not neighbor-optimal.** In this scenario, we will prove that $\bar{p}(R)$ is finite to obtain $\lim_{R \rightarrow \infty} \bar{p}(R)/R = 0 \neq q$. To start, assume that conic neighborhood ϵ of q is small enough so that selected-tokens α remain unchanged within C_ϵ . This is without generality because if directional convergence fails in a small neighborhood of q , it will fail in the larger neighborhood as well. Secondly, if $\lim_{R \rightarrow \infty} \|\bar{p}(R)\| \rightarrow \infty$ and $\bar{p}(R) \in C_\epsilon$, since softmax will eventually perfectly select α (i.e. assigning probability 1 on token indices (i, α_i)), we would have

$$\lim_{R \rightarrow \infty} \mathcal{L}(\bar{p}(R)) = \mathcal{L}_\star = \frac{1}{n} \sum_{i=1}^n \ell(\gamma_{i\alpha_i}).$$

Note that, this is simply by selection of α and regardless of $\bar{p}(R)$ directionally converges to q . This means that, if there exists a finite $p \in C_\epsilon$ such that $\mathcal{L}(p) < \mathcal{L}_\star$ (i.e. outperforming the training loss of $\|\bar{p}(R)\| \rightarrow \infty$), then $\|\bar{p}(R)\| < \infty$. This would conclude the proof.

Thus, we will simply find such a p obeying $\mathcal{L}(p) < \mathcal{L}_\star$. To this aim, we first prove the following lemma.

Lemma 9 *Given a fixed unit Euclidean norm vector p , if all directional neighbors of p consistently have higher scores for their associated selected tokens, i.e., $\gamma_{i\alpha_i} < \gamma_{i\beta}$ for all (i, α_i) and directional neighbor (i, β) , then there exists $\bar{R}$ such that for all $R > \bar{R}$,*

$$\mathcal{L}(R \cdot p) < \mathcal{L}_\star = \lim_{R \rightarrow \infty} \mathcal{L}(R \cdot p). \quad (69)$$

Proof. Define the maximum score difference $K_+ = \sup_{i \in [n], t \neq \alpha_i} |\gamma_{i\alpha_i} - \gamma_{it}|$. Also let $\mathcal{M}_p$ be the set of directional neighbors achieving the minimum margin Γ_p ; see (61). Define $\Gamma^{it} = k_{i\alpha_i}^\top p - k_{it}^\top p$. Define δ_p to be the margin difference between the directional-neighbors and the second-most minimum-margin neighbors defined as

$$\delta_p = \min_{i \in [n], t \neq \alpha_i, (i,t) \notin \mathcal{M}_p} \Gamma_p^{it} - \Gamma_p. \quad (70)$$

⁶Here, we apply this lemma to compare q' against all $p \in \mathcal{N}_\delta$. We can do this uniform comparison because the R requirement in Lemma 8 only depends on the margin difference and global problem variables and not the particular choice of $p \in \mathcal{N}_\delta$.

To proceed, setting $\kappa = \min_{(j,\beta) \in \mathcal{M}_p} \gamma_{j\beta} - \gamma_{j\alpha_j} > 0$ and using (63), we can bound

$$\begin{aligned} \mathcal{L}(Rp) - \mathcal{L}_\star &\leq -\frac{1}{n} \sum_{(j,\beta) \in \mathcal{M}_p} M_{j\beta} e^{-R\Gamma_p \kappa} + \frac{1}{n} \sum_{i \in [n], t \neq \alpha_i, (i,t) \notin \mathcal{M}_p} M_{it} e^{-R\Gamma_p^i} (\gamma_{i\alpha_i} - \gamma_{it}) \\ &\leq -\frac{1}{n} M_- e^{-R\Gamma_p \kappa} + M_+ K_+ T e^{-R(\Gamma_p + \delta_p)}. \end{aligned}$$

Consequently, we have found that $\mathcal{L}(Rp) < \mathcal{L}_\star$ as soon as

$$\frac{1}{n} M_- e^{-R\Gamma_p \kappa} \geq M_+ K_+ T e^{-R(\Gamma_p + \delta_p)}.$$

This happens when $R \gtrsim 1/\delta_q$ (up to logarithmic terms) establishing the desired statement. $\blacksquare$

Based on this lemma, what we need is constructing a perturbation to modify $\mathbf{q}$'s directional neighbors and make sure all of them have strictly better scores than their associated selected-tokens. Note that, once we construct a new candidate (say $\mathbf{q}_0 = \mathbf{q} + \text{perturbation}$), all sufficiently large scalings of $\mathbf{q}_0$ will achieve $\mathcal{L}(R \cdot \mathbf{q}_0) < \mathcal{L}_\star$. Thus, we can find a strictly better solution than $\mathcal{L}_\star$ for any norm lower bound R_0 – which is enforced within the definition of C_ϵ .

Lemma 10 *There are at most d directional neighbors i.e. $|\mathcal{M}_q| \leq d$.*

Proof. Directional neighbors are indices (i, t) obeying the inequality

$$(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{q} = \Gamma_q.$$

Declare $\mathbf{D}$ to be the matrix with rows obtained by these key differences $\mathbf{k}_{it} - \mathbf{k}_{i\alpha_i}$. $\mathbf{D} \in \mathbb{R}^{M \times d}$ where $M = |\mathcal{M}_q|$. We then obtain $\mathbf{D}\mathbf{q} = -\Gamma_q \mathbf{1}_M$ where $\mathbf{1}_M$ is the all ones vector. If $M > d$, the equality $\mathbf{D}\mathbf{q} = -\Gamma_q \mathbf{1}_M$ cannot be satisfied because $\mathbf{1}_M$ is not in the range space of $\mathbf{D}$ by our assumption of general key embedding positions. $\blacksquare$

To proceed, $|\mathcal{M}_q| \leq d$ and let $\mathbf{D}$ be as defined in the lemma above. $\mathbf{D}$ is also full-rank by our assumption of general key positions. We use $\mathbf{D}$ to construct a perturbation as follows. Let $\mathcal{M}_q^+ \subset \mathcal{M}_q$ be the set of directional neighbor with strictly higher scores than their associated selected-tokens. In other words, all $(j, \beta) \in \mathcal{M}_q^+$ obeys

$$\gamma_{j\beta} > \gamma_{j\alpha_j}.$$

Define the score difference $\kappa = \min_{(j,\beta) \in \mathcal{M}_q^+} \gamma_{j\beta} - \gamma_{j\alpha_j} > 0$. We know $\kappa > 0$ because α is not neighbor-optimal. Finally, define the indicator vector of $\mathbf{1}_+$ with same dimension as cardinality $|\mathcal{M}_q|$. $\mathbf{1}_+$ is 1 for the rows of $\mathbf{D}$ corresponding to $\mathcal{M}_q^+$ and is 0 otherwise. Finally, set the perturbation as

$$\mathbf{q}^\perp = \mathbf{D}^\dagger \mathbf{1}_+.$$

where we used the full-rankness of $\mathbf{D}$ during pseudo-inversion. To proceed, for a small $\epsilon_0 > 0$, consider the candidate direction $\mathbf{q}_0 = \mathbf{q} + \epsilon_0 \mathbf{q}^\perp$. We pick $\epsilon_0 = O(\epsilon)$ sufficiently small to ensure $\mathbf{q}_0 \in C_\epsilon$. To finalize, let us consider the margins of the tokens within $\mathbf{q}_0$. Similar to Lemma 9, set $\delta_q = \min_{i \in [n], t \neq \alpha_i, (i,t) \notin \mathcal{M}_q} \Gamma_q^{it} - \Gamma_q > 0$. Let $\bar{\epsilon}_0 = \|\mathbf{q}_0\| - 1 = \|\mathbf{q} + \epsilon_0 \mathbf{q}^\perp\| - 1$. Using definition of $\mathbf{q}^\perp$, we have that

- For $(i, t) \in \mathcal{M}_q^+$, we achieve a margin of

$$(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top (\mathbf{q} + \epsilon_0 \mathbf{q}^\perp) / (1 + \bar{\epsilon}_0) = \frac{\Gamma_q - \epsilon_0}{1 + \bar{\epsilon}_0}.$$

- For $(i, t) \in \mathcal{M}_q^+$, we achieve a margin of $\frac{\Gamma_q}{1 + \bar{\epsilon}_0}$.
- For $(i, t) \notin \mathcal{M}_q$, setting $K = \|\mathbf{q}^\perp\| \cdot \sup_{i,t,\tau} \|\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}\|$, we achieve a margin of at most

$$(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top (\mathbf{q} + \epsilon_0 \mathbf{q}^\perp) / (1 + \bar{\epsilon}_0) \geq \frac{\Gamma_q + \delta_q - \epsilon_0 K}{1 + \bar{\epsilon}_0}.$$

In short, since $\bar{\epsilon}_0 = O(\epsilon_0)$, setting ϵ_0 sufficiently small guarantees that $\mathcal{M}_q^+$ is the set of directional neighbors of q_0 . Since $\mathcal{M}_q^+$ has strictly higher scores than their associated selected-tokens, applying Lemma 9 on q_0 shows that, $\mathcal{L}(R \cdot q_0) < \mathcal{L}_\star$ for sufficiently large R implying $\|\tilde{p}(R)\| < \infty$.

• **(B) q selects multiple tokens for some inputs $i \in [n]$:** In this setting, we will again construct a perturbation to create a scenario where $q_0 = q + \epsilon q^\perp$ selects a single token for each input $i \in [n]$. We will then employ margin analysis (first statement of Lemma 8) to conclude that q_0 outperforms a $\delta \ll \epsilon$ neighborhood of q .

Let $\mathcal{I} \subset [n]$ be the set of inputs for which q selects multiple tokens. Specifically, for each $i \in \mathcal{I}$, there is $\mathcal{T}_i \subset [T]$ such that $|\mathcal{T}_i| \geq 2$ and for any $i \in \mathcal{I}$ and $\theta \in \mathcal{T}_i$,

$$k_{i\theta}^\top q = \arg \max_{t \in [T]} k_{it}^\top q.$$

From these multiply-selected token indices let us select the highest score one, namely, $\beta_i = \arg \max_{\theta \in \mathcal{T}_i} \gamma_{i\theta}$ for $i \in \mathcal{I}$. Now, define the unique optimal tokens for each input as $\alpha \in \mathbb{R}^n$ where $\alpha_i := \beta_i$ for $i \in \mathcal{I}$ and $\alpha_i = \arg \max_{t \in [T]} k_{it}^\top q$ for $i \notin \mathcal{I}$. Define $\mathcal{L}_\star = \frac{1}{n} \sum_{i=1}^n \ell(\gamma_{i\alpha_i})$ as earlier.

Secondly, we construct a perturbation $q^\perp$ to show that $q_0 = q + \epsilon_0 q^\perp$ can select tokens α asymptotically. To see this, define the matrix D where each (unique) row is given by $k_{i\alpha_i} - k_{i\theta}$ where $\theta \in \mathcal{T}_i, \theta \neq \alpha_i, i \in \mathcal{I}$. Now note that, $(k_{i\alpha_i} - k_{i\theta})^\top q = 0$ for all $\theta \in \mathcal{T}_i, \theta \neq \alpha_i, i \in \mathcal{I}$. Since keys are in general positions, this implies that D has at most $d - 1$ rows and, thus, its rows are linearly independent. Consequently, choose $q^\perp = D^\dagger \mathbf{1}$ where $\dagger$ denotes pseudo-inverse and $\mathbf{1}$ is the all ones vector. Also let Γ_q be the margin of directional margin of q that is

$$\Gamma_q = \min_{i \in [n], t \notin \mathcal{T}_i} (k_{i\alpha_i} - k_{it})^\top q.$$

With this choice and setting $K = \sup_{i,t,\tau} \|k_{i\alpha_i} - k_{it}\|$, we have that

- For $\theta \in \mathcal{T}_i, \theta \neq \alpha_i$: $(k_{i\alpha_i} - k_{i\theta})^\top q_0 = (k_{i\alpha_i} - k_{i\theta})^\top q^\perp = \epsilon_0$.
- For all other (i, t) with $t \neq \alpha_i$: $(k_{i\alpha_i} - k_{it})^\top q_0 \geq \Gamma_q - K \|q^\perp\| \epsilon_0$.

Choosing $\epsilon_0 < \Gamma_q / (1 + K \|q^\perp\|)$, together, these imply that,

- $\alpha = (\alpha_i)_{i=1}^n$ is the selected-tokens of q_0 ,
- q_0 achieves a directional margin of

$$\Gamma_{q_0} = \frac{\epsilon_0}{\|q_0\|} \geq \frac{\epsilon_0}{1 + \epsilon_0 \|q^\perp\|},$$

- q_0 is neighbor-optimal because directional neighbors are $\theta \in \mathcal{T}_i, \theta \neq \alpha_i$ and $\gamma_{i\theta} < \gamma_{i\alpha_i}$.

Note that, these conditions lay the groundwork for applying the first statement of Lemma 8 with $p \leftarrow q_0$. We next explore the optimal directions within $\mathcal{N}_\delta$ and show that q_0 strictly outperform them in terms of training loss.

To proceed, given small $\delta \ll \epsilon_0$, let us study $\mathcal{L}_R = \min_{p \in \mathcal{N}_\delta, R \leq \|p\| < \infty} \mathcal{L}(p)$. Here, recall that p has a δ -small directional perturbation around q which can modify the token selections by breaking the ties between the multiply-selected token indices by q . However, thanks to the distinct token score assumption, for large $R \leq \|p\|$, the optimal $p \in \mathcal{N}_\delta$ is guaranteed to (uniquely) select indices $\alpha_i \in \mathcal{T}_i$. Because all other p directions – which, asymptotically, either select other tokens $\theta \in \mathcal{T}_i, \theta \neq \alpha_i$ or split the probabilities equally across a subset of $\mathcal{T}_i$ – achieve a larger loss. For instance, q will split the probabilities equally across $\mathcal{T}_i$ to achieve an asymptotic loss of

$$\mathcal{L}_\star^q := \lim_{R \rightarrow \infty} \mathcal{L}(R \cdot q) = \frac{1}{n} \sum_{i \notin \mathcal{I}} \ell(\gamma_{i\alpha_i}) + \frac{1}{n} \sum_{i \in \mathcal{I}} \ell\left(\frac{1}{|\mathcal{T}_i|} \sum_{\theta \in \mathcal{T}_i} \gamma_{i\theta}\right)$$

Thus, $\mathcal{L}_\star^q > \mathcal{L}_\star$ because $\ell\left(\frac{1}{|\mathcal{T}_i|} \sum_{\theta \in \mathcal{T}_i} \gamma_{i\theta}\right) > \ell(\gamma_{i\alpha_i}) = \ell(\gamma_{i\alpha_i})$ where α_i has the highest score i.e. $\gamma_{i\alpha_i} > \frac{1}{|\mathcal{T}_i|} \sum_{\theta \in \mathcal{T}_i} \gamma_{i\theta}$. Set $\tilde{p}(R) = \arg \min_{p \in \mathcal{N}_\delta, \|p\| \leq R} \mathcal{L}(p)$. Consequently, there are two scenarios are:

- $\lim_{R \rightarrow \infty} \|\tilde{p}(R)\|$ is finite. This already proves the statement of the theorem as $\tilde{p}(R)/R \rightarrow 0$ within $\delta < \epsilon$ neighborhood of q .

- For sufficiently large R , the selected-tokens of $\tilde{\mathbf{p}}(R)$ are $\alpha = (\alpha_i)_{i=1}^n$.

Proceeding with the second (remaining scenario), we study the directional margin of $\tilde{\mathbf{p}}(R)$. More broadly, for any $\mathbf{p} \in \mathcal{N}_\delta$ and $\tilde{\mathbf{p}} = \mathbf{p}/\|\mathbf{p}\|$ with selected-tokens α , using the fact that $\|\tilde{\mathbf{p}} - \mathbf{q}\| \leq \delta \ll \epsilon_0$, we can bound the directional margin as

- For $\theta \in \mathcal{T}_i, \theta \neq \alpha_i$: $(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{i\theta})^\top \tilde{\mathbf{p}} = (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{i\theta})^\top (\tilde{\mathbf{p}} - \mathbf{q}) \leq K\delta$.
- For all other (i, t) with $t \neq \alpha_i$: $(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \tilde{\mathbf{p}} \geq \Gamma_q - K\delta$.

This means that, any such $\tilde{\mathbf{p}} \in \mathcal{N}_\delta$ achieves a directional margin of at most

$$\Gamma_{\tilde{\mathbf{p}}} \leq K\delta.$$

Applying Lemma 8 and setting $\delta = O(\epsilon_0)$, this implies that for

$$R \gtrsim \frac{1}{\Gamma_{q_0} - \Gamma_{\tilde{\mathbf{p}}}} = \frac{1}{\frac{\epsilon_0}{1+\epsilon_0\|\mathbf{q}^\perp\|} - K\delta} = O\left(\frac{1}{\epsilon_0}\right),$$

we have that $\mathcal{L}(R \cdot \mathbf{q}_0) < \min_{\|\mathbf{p}\|=R, \mathbf{p} \in \mathcal{N}_\delta} \mathcal{L}(\mathbf{p})$. Since this holds for all R , (68) holds (similar to Case (A1)) and we conclude that whenever $\|\tilde{\mathbf{p}}(R)\| \rightarrow \infty$, it doesn't directionally converge within $\mathcal{N}_\delta$ (i.e. $\delta > 0$ neighborhood of $\mathbf{q}$) proving the advertised result. ■

B.6 Proof of Lemma 2

We prove a slightly general restatement where we require $\mathbf{v} \in \text{range}(\mathbf{W}^\top)$ – instead of full-rank $\mathbf{W}$.

Lemma 11 Suppose for all $i \in [n]$ and $t \neq \text{opt}_i$, $Y_i = 1$ and $\gamma_{it} < \gamma_{i\text{opt}_i}$. Also suppose $\mathbf{v} \in \text{range}(\mathbf{W}^\top)$. Then, $\mathbf{p}^{\text{mm}\star}$ exists – i.e. (ATT-SVM) is feasible for optimal indices $\alpha_i \leftarrow \text{opt}_i$.

Proof. To establish the existence of $\mathbf{p}^{\text{mm}\star}$, we only need to find a direction that demonstrates the feasibility of (ATT-SVM), i.e. we need to find $\mathbf{p}$ that satisfies the margin constraints. To begin, let's define the minimum score difference:

$$\underline{\gamma} = \min_{i \in [n], t \neq \text{opt}_i} \gamma_{i\text{opt}_i} - \gamma_{it}.$$

We then set $\mathbf{p} = \underline{\gamma}^{-1}(\mathbf{W}^\top)^\dagger \mathbf{v}$ where † denotes pseudo-inverse. By assumption $\mathbf{W}^\top \mathbf{p} = \underline{\gamma}^{-1} \mathbf{v}$. To conclude, observe that $\mathbf{p}$ is a feasible solution since $\mathbf{k}_{it} = \mathbf{W} \mathbf{x}_{it}$ and for all $i \in [n]$ and $t \neq \text{opt}_i$, we have that

$$\begin{aligned} (\mathbf{k}_{i\text{opt}_i} - \mathbf{k}_{it})^\top \mathbf{p} &= (\mathbf{x}_{i\text{opt}_i} - \mathbf{x}_{it})^\top \mathbf{W}^\top \mathbf{p} = \underline{\gamma}^{-1} (\mathbf{x}_{i\text{opt}_i} - \mathbf{x}_{it})^\top \mathbf{W}^\top (\mathbf{W}^\top)^\dagger \mathbf{v} \\ &= \underline{\gamma}^{-1} (\mathbf{x}_{i\text{opt}_i} - \mathbf{x}_{it})^\top \mathbf{v} \geq 1, \end{aligned}$$

which together with the constraints in (ATT-SVM) completes the proof. ■

C Addendum to Section 3

C.1 Proof of Theorem 5

Proof. Suppose the claim is incorrect and either $\mathbf{p}_R/R$ or $\mathbf{v}_r/r$ fails to converge as R, r grows. Set $\Xi = 1/\|\mathbf{p}^{\text{mm}}\|$, $\Gamma = 1/\|\mathbf{v}^{\text{mm}}\|$, $\tilde{\mathbf{p}}^{\text{mm}} = R\Xi \mathbf{p}^{\text{mm}}$ and $\tilde{\mathbf{v}}^{\text{mm}} = r\Gamma \mathbf{v}^{\text{mm}}$. The proof strategy is obtaining a contradiction by proving that $(\tilde{\mathbf{v}}^{\text{mm}}, \tilde{\mathbf{p}}^{\text{mm}})$ is a strictly better solution compared to $(\mathbf{v}_r, \mathbf{p}_R)$ for large R, r . Without losing generality, we will set $\alpha_i = 1$ for all $i \in [n]$ as the problem is invariant to tokens' permutation. Define $q_i^p = 1 - s_{i1}^p$ to be the amount of non-optimality (cumulative probability of non-first tokens) where $s_i^p = \mathbb{S}(\mathbf{K}_i \mathbf{p})$ is the softmax probabilities.

• **Case 1: $\mathbf{p}_R/R$ does not converge.** Under this scenario there exists $\delta, \gamma = \gamma(\delta) > 0$ such that we can find arbitrarily large R with $\|\mathbf{p}_R/R - \tilde{\mathbf{p}}^{\text{mm}}/R\| \geq \delta$ and margin induced by $\mathbf{p}_R/R$ is at most $\Xi(1 - \gamma)$ (from strong convexity of (ATT-SVM)). Following q_i^p definition above, set $\hat{q}_{\max} = \sup_{i \in [n]} q_i^{\mathbf{p}_R}$ to be

worst non-optimality in $\mathbf{p}_R$ and $q_{\max}^* = \sup_{i \in [n]} q_i^{\tilde{\mathbf{p}}^{mm}}$ to be the same for $\tilde{\mathbf{p}}^{mm}$. Repeating the identical argument in Theorem 8 (specifically (84)), we can bound the non-optimality amount $q_i^{\tilde{\mathbf{p}}^{mm}}$ of $\tilde{\mathbf{p}}^{mm}$ as

$$q_i^{\tilde{\mathbf{p}}^{mm}} = \frac{\sum_{t \neq \alpha_i} \exp(\mathbf{k}_{it}^\top \tilde{\mathbf{p}}^{mm})}{\sum_{t \in [T]} \exp(\mathbf{k}_{it}^\top \tilde{\mathbf{p}}^{mm})} \leq \frac{\sum_{t \neq \alpha_i} \exp(\mathbf{k}_{it}^\top \tilde{\mathbf{p}}^{mm})}{\exp(\mathbf{k}_{i\alpha_i}^\top \tilde{\mathbf{p}}^{mm})} \leq T \exp(-R\Xi). \quad (71)$$

Thus, $q_{\max}^* = \max_{i \in [n]} q_i^{\tilde{\mathbf{p}}^{mm}} \leq T \exp(-R\Xi)$. Next without losing generality, assume first margin constraint is γ -violated by $\mathbf{p}_R$ and $\min_{t \neq \alpha_1} (\mathbf{k}_{1\alpha_1} - \mathbf{k}_{1t})^\top \mathbf{p}_R \leq \Xi R(1 - \gamma)$. Denoting the amount of non-optimality of the first input as $q_1^{\mathbf{p}_R}$, we find

$$q_1^{\mathbf{p}_R} = \frac{\sum_{t \neq \alpha_1} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)}{\sum_{t \in [T]} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)} \geq \frac{1}{T} \frac{\sum_{t \neq \alpha_1} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)}{\exp(\mathbf{k}_{1\alpha_1}^\top \mathbf{p}_R)} \geq T^{-1} \exp(-(1 - \gamma)R\Xi). \quad (72)$$

We similarly have $q_{\max}^* \geq T^{-1} \exp(-R\Xi)$ to find that

$$\begin{aligned} \log(\hat{q}_{\max}) &\geq -(1 - \gamma)\Xi R - \log T, \\ -\Xi R - \log T &\leq \log(q_{\max}^*) \leq -\Xi R + \log T. \end{aligned} \quad (73)$$

In words, $\tilde{\mathbf{p}}^{mm}$ contains exponentially less non-optimality compared to $\mathbf{p}_R$ as R grows. The remainder of the proof differs from Theorem 8 as we need to upper/lower bound the logistic loss of $(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{mm})$ and $(\mathbf{v}_r, \mathbf{p}_R)$ respectively to conclude with the contradiction.

First, let us upper bound the logistic loss of $(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{mm})$. Set $\mathbf{r}_i = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \tilde{\mathbf{p}}^{mm})$. Observe that if $\|\mathbf{r}_i - \mathbf{x}_{i1}\| \leq \epsilon_i$, we have that $\tilde{\mathbf{v}}^{mm}$ satisfies the SVM constraints on $\mathbf{r}_i$ with $Y_i \cdot \mathbf{r}_i^\top \tilde{\mathbf{v}}^{mm} \geq 1 - \epsilon_i/\Gamma$. Consequently, setting $\epsilon_{\max} = \sup_{i \in [n]} \epsilon_i$, $\tilde{\mathbf{v}}^{mm}$ achieves a label-margin of $\Gamma - \epsilon_{\max}$ on the dataset $(Y_i, \mathbf{r}_i)_{i \in [n]}$. With this, we upper bound the logistic loss of $(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{mm})$ as follows. Let $M = \sup_{i \in [n], t, \tau \in [T]} \|\mathbf{x}_{it} - \mathbf{x}_{i\tau}\|$. Let us recall the fact (73) that worst-case perturbation is

$$\epsilon_{\max} \leq M \exp(-\Xi R + \log T) = MT \exp(-\Xi R).$$

This implies that

$$\begin{aligned} \mathcal{L}(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{mm}) &\leq \max_{i \in [n]} \log(1 + \exp(-Y_i \mathbf{r}_i^\top \tilde{\mathbf{v}}^{mm})) \\ &\leq \max_{i \in [n]} \exp(-Y_i \mathbf{r}_i^\top \tilde{\mathbf{v}}^{mm}) \\ &\leq \exp(-r\Gamma + r\epsilon_{\max}) \\ &\leq e^{rMT \exp(-\Xi R)} e^{-r\Gamma}. \end{aligned} \quad (74)$$

Conversely, we obtain a lower bound for $(\mathbf{v}_r, \mathbf{p}_R)$. Set $\mathbf{r}_i = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}_R)$. Using Assumption C, we find that solving (SVM) on $(Y_i, \mathbf{r}_i)_{i \in [n]}$ achieves at most $\Gamma - \nu e^{-(1-\gamma)\Xi R}/T$ margin. Consequently, we have

$$\begin{aligned} \mathcal{L}(\mathbf{v}_r, \mathbf{p}_R) &\geq \frac{1}{n} \max_{i \in [n]} \log(1 + \exp(-Y_i \mathbf{r}_i^\top \mathbf{v}_r)) \\ &\geq \frac{1}{2n} \max_{i \in [n]} \exp(-Y_i \mathbf{r}_i^\top \mathbf{v}_r) \wedge \log 2 \\ &\geq \frac{1}{2n} \exp(-r(\Gamma - \nu e^{-(1-\gamma)\Xi R}/T)) \wedge \log 2 \\ &\geq \frac{1}{2n} e^{r(\nu/T) \exp(-(1-\gamma)\Xi R)} e^{-r\Gamma} \wedge \log 2. \end{aligned} \quad (75)$$

Observe that, this lower bound dominates the previous upper bound when R is large, namely, when (ignoring the multiplier $1/2n$ for brevity)

$$(\nu/T) e^{-(1-\gamma)\Xi R} \geq MT e^{-\Xi R} \iff R \geq R_0 := \frac{1}{\gamma\Xi} \log\left(\frac{MT^2}{\nu}\right).$$

Thus, we indeed obtain the desired contradiction since such large R is guaranteed to exist when $\mathbf{p}_R/R \rightarrow \mathbf{p}^{mm}$.

• **Case 2: $\mathbf{v}_r/r$ does not converge.** This is the simpler scenario: There exists $\delta > 0$ such that we can find arbitrarily large r obeying $\|\mathbf{v}_r/r - \mathbf{v}^{mm}\|/\|\mathbf{v}^{mm}\| \geq \delta$. If $\|\mathbf{p}_R/R - \Xi \mathbf{p}^{mm}\| \rightarrow 0$, then “Case

1" applies. Otherwise, we have $\|\mathbf{p}_R/R - \Xi \mathbf{p}^{mm}\| \rightarrow 0$, thus we can assume $\|\mathbf{p}_R/R - \Xi \mathbf{p}^{mm}\| \leq \epsilon$ for arbitrary choice of $\epsilon > 0$.

On the other hand, due to the strong convexity of (SVM), for some $\gamma := \gamma(\delta) > 0$, $\mathbf{v}_r$ achieves a margin of at most $(1 - \gamma)\Gamma r$ on the dataset $(Y_i, \mathbf{x}_{i1})_{i \in [n]}$. Additionally, since $\|\mathbf{p}_R/R - \Xi \mathbf{p}^{mm}\| \leq \epsilon$, $\mathbf{p}_R$ strictly separates all optimal tokens (for small enough $\epsilon > 0$) and $\hat{q}_{\max} := f(\epsilon) \rightarrow 0$ as $R \rightarrow \infty$. Consequently, setting $\mathbf{r}_i = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{p}_R)$, for sufficiently large $R > 0$ setting $M = \sup_{i \in [n], t \in [T]} \|\mathbf{x}_{it}\|$, we have that

$$\begin{aligned} \min_{i \in [n]} Y_i \mathbf{v}_r^\top \mathbf{r}_i &\leq \min_{i \in [n]} Y_i \mathbf{v}_r^\top \mathbf{x}_{i1} + \sup_{i \in [n]} |\mathbf{v}_r^\top (\mathbf{r}_i - \mathbf{x}_{i1})| \\ &\leq (1 - \gamma)\Gamma r + M f(\epsilon) r \\ &\leq (1 - \gamma/2)\Gamma r. \end{aligned} \tag{76}$$

This in turn implies that logistic loss is lower bounded by (following (75)),

$$\mathcal{L}(\mathbf{v}_r, \mathbf{p}_R) \geq \frac{1}{2n} e^{\gamma \Gamma r/2} e^{-\Gamma r} \wedge \log 2.$$

Going back to (74), this exponentially dominates the upper bound of $(\tilde{\mathbf{p}}^{mm}, \tilde{\mathbf{v}}^{mm})$ whenever $rMT \exp(-\Xi R) < r\gamma\Gamma/2$, (that is, whenever R, r are sufficiently large), again concluding the proof. ■

C.2 Proof of Theorem 6

We will prove this result in two steps. Our first claim restricts the optimization to the particular quadrant induced by $\min_{t \neq \alpha_i} (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{p}_R \geq 0$ under the theorem's condition $\mathbb{S}(\mathbf{K}_i \mathbf{p}_R)_{\alpha_i} \rightarrow 1$.

Lemma 12 Suppose $\mathbb{S}(\mathbf{K}_i \mathbf{p}_R)_{\alpha_i} \rightarrow 1$. Then, there exists R_0 such that for all $R \geq R_0$, we have that,

$$\min_{t \neq \alpha_i} (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{p}_R \geq 0, \quad \text{for all } i \in [n]. \tag{77}$$

Proof. Suppose the claim does not hold. Set $\mathbf{s}_i^R = \mathbb{S}(\mathbf{K}_i \mathbf{p}_R)$. Fix R_0 such that $\mathbf{s}_{i\alpha_i}^R \geq 0.9$ for all $R \geq R_0$. On the other hand, there exists arbitrarily large R for which $(\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it})^\top \mathbf{p}_R < 0$ for some $t \neq \alpha_i \in [T]$ and $i \in [n]$. At this (R, i, t) choices, we have that $\mathbf{s}_{it}^R \geq \mathbf{s}_{i\alpha_i}^R$. Since $\mathbf{s}_{it}^R + \mathbf{s}_{i\alpha_i}^R \leq 1$, we find $\mathbf{s}_{i\alpha_i}^R < 0.5$ which contradicts with $\mathbf{s}_{i\alpha_i}^R \geq 0.9$. ■

Let $\mathcal{Q}$ be the set of $\mathbf{p}$ satisfying the quadrant constraint (77) – i.e. indices $(\alpha_i)_{i=1}^n$ are selected. Let $\mathbf{h}_R$ be the solution of regularization path of $(\mathbf{v}, \mathbf{p})$ subject to the constraint $\mathbf{p} \in \mathcal{Q}$. From Lemma 12, we know that, for some R_0 and all $R \geq R_0$, $\mathbf{h}_R = \mathbf{p}_R$. Thus, if the limit exists, we have that $\lim_{R \rightarrow \infty} \mathbf{h}_R/R = \lim_{R \rightarrow \infty} \mathbf{p}_R/R$.

To proceed, we will prove that $\lim_{R \rightarrow \infty} \mathbf{h}_R/R$ exists and is equal to $\mathbf{p}^{relax}/\|\mathbf{p}^{relax}\|$ and simultaneously establish $\mathbf{v}_r/r \rightarrow \mathbf{v}^{mm}/\|\mathbf{v}^{mm}\|$.

Lemma 13 $\lim_{R \rightarrow \infty} \mathbf{h}_R/R = \mathbf{p}^{relax}/\|\mathbf{p}^{relax}\|$ and $\lim_{r \rightarrow \infty} \mathbf{v}_r/r = \mathbf{v}^{mm}/\|\mathbf{v}^{mm}\|$.

Proof. The proof will be similar to that of Theorem 5. As usual, we aim to show that SVM-solutions constitute the most competitive direction. Set $\Xi = 1/\|\mathbf{p}^{relax}\|$.

• **Case 1: $\mathbf{h}_R/R$ does not converge.** Under this scenario there exists $\delta, \gamma = \gamma(\delta) > 0$ such that we can find arbitrarily large R with $\|\mathbf{h}_R/R - \Xi \mathbf{p}^{relax}\| \geq \delta$. This implies that margin induced by $\mathbf{h}_R/R$ is at most $\Xi(1 - \gamma)$ over the support vectors $\mathcal{S}$ (from strong convexity of (10)). The reason is that, $\mathbf{h}_R$ satisfies $\mathbf{h}_R^\top (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}) \geq 0$ for all $t \neq \alpha_i$ by construction as $\mathbf{h}_R \in \mathcal{Q}$. Thus, a constraint over the support vectors have to be violated (when normalized to the same ℓ_2 norm as $\|\mathbf{p}^{relax}\| = 1/\Xi$).

As usual, we will construct a solution strictly superior to $\mathbf{h}_R$ and contradicts with its optimality.

Construction of competitor: Rather than using $\mathbf{p}^{relax}$ direction, we will choose a slightly deviating direction that ensures the selection of the correct tokens over non-supports $\tilde{\mathcal{S}}$. Specifically, consider the solution of (10) where we tighten the non-support constraints by arbitrarily small $\epsilon > 0$.

$$\mathbf{p}^{\epsilon-relx} = \arg \min_{\mathbf{p}} \|\mathbf{p}\| \quad \text{such that} \quad \mathbf{p}^\top (\mathbf{k}_{i\alpha_i} - \mathbf{k}_{it}) \geq \begin{cases} 1 & \text{for all } t \neq \alpha_i, i \in \mathcal{S} \\ \epsilon & \text{for all } t \neq \alpha_i, i \in \tilde{\mathcal{S}} \end{cases}. \tag{78}$$

Let $\mathbf{p}^{mm}$ be the solution of (ATT-SVM) with $\alpha = (\alpha_i)_{i=1}^n$ (which was assumed to be separable). Observe that $\mathbf{p}_\epsilon^{mm} = \epsilon \mathbf{p}^{mm} + (1 - \epsilon) \mathbf{p}^{relax}$ satisfies the constraints of (78). Additionally, $\mathbf{p}_\epsilon^{mm}$ would achieve a margin of $\frac{1}{(1-\epsilon)/\Xi + \epsilon/\Delta} = \frac{\Delta\Xi}{\Delta + \epsilon(\Xi - \Delta)}$ where $\Delta = 1/\|\mathbf{p}^{mm}\|$. Using optimality of $\mathbf{p}^{\epsilon-rlx}$, this implies that the reduced margin $\Xi_\epsilon = 1/\|\mathbf{p}_\epsilon^{\epsilon-rlx}\|$ (by enforcing ϵ over non-support) over the support vectors is a Lipschitz function of ϵ . That is $\Xi_\epsilon \geq \Xi - \epsilon M$ for some $M \geq 0$. To proceed, choose an $\epsilon > 0$ such that, it is strictly superior to margin induced by $\mathbf{h}_R$, that is,

$$\Xi_\epsilon \geq \Xi \left(1 - \frac{\gamma}{2}\right).$$

To proceed, set $\tilde{\mathbf{p}}^{\epsilon-rlx} = R\Xi_\epsilon \mathbf{p}^{\epsilon-rlx}$. Let us recall the following notation from the proof of Theorem 5: $s_i^p = \mathbb{S}(\mathbf{K}_i \mathbf{p})$ and $q_i^p = 1 - s_{i\alpha_i}$. Set $\hat{q}_{\max} = \max_{i \in \mathcal{S}} q_i^{\mathbf{h}_R}$ to be worst non-optimality of $\mathbf{h}_R$ over **support set**. Similarly, define $q_{\max}^* = \max_{i \in \mathcal{S}} q_i^{\tilde{\mathbf{p}}^{\epsilon-rlx}}$ to be the same for $\tilde{\mathbf{p}}^{\epsilon-rlx}$. Repeating the identical arguments to (71), (72), (73), and using the fact that $\mathbf{p}^{\epsilon-rlx}$ achieves a margin $\Xi(1 - \frac{\gamma}{2}) \leq \Xi_\epsilon \leq \Xi$, we end up with the lines

$$\log(\hat{q}_{\max}) \geq -(1 - \gamma)\Xi R - \log T, \quad (79a)$$

$$-\Xi R - \log T \leq \log(q_{\max}^*) \leq -\Xi(1 - 0.5\gamma)R + \log T. \quad (79b)$$

In what follows, we will prove that $\tilde{\mathbf{p}}^{\epsilon-rlx}$ achieves a strictly smaller logistic loss contradicting with the optimality of $\mathbf{p}_R$ (whenever $\|\mathbf{h}_R/R - \Xi \mathbf{p}^{relax}\| \geq \delta$).

Upper bounding logistic loss. Let us now upper bound the logistic loss of $(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{\epsilon-rlx})$ where $\tilde{\mathbf{v}}^{mm} = r\Gamma \mathbf{v}^{mm}$ with $\mathbf{v}^{mm}$ being the solution of (SVM) with $\mathbf{r}_i \leftarrow \mathbf{x}_{i\alpha_i}$ and $\Gamma = 1/\|\mathbf{v}^{mm}\|$. Set $\mathbf{r}_i = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \tilde{\mathbf{p}}^{\epsilon-rlx})$. Set $\nu = \min_{i \in \bar{\mathcal{S}}} Y_i \cdot \mathbf{x}_{i\alpha_i}^\top \mathbf{v}^{mm} - 1$ to be the additional margin buffer that non-support vectors have access to. Also set $M = \sup_{i \in [n], t, \tau \in [T]} \|\mathbf{x}_{it} - \mathbf{x}_{i\tau}\|$. Observe that we can write

$$\mathbf{x}_{i\alpha_i} - \mathbf{r}_i = \sum_{t \neq \alpha_i} s_{it}(\mathbf{x}_{i\alpha_i} - \mathbf{x}_{it}) \implies \|\mathbf{x}_{i\alpha_i} - \mathbf{r}_i\| \leq q_i M.$$

Non-supports achieve strong label-margin: Using above and (78) for all $i \in \bar{\mathcal{S}}$ and $t \neq \alpha_i$, we have that $s_{it} \leq e^{-\epsilon\Xi_\epsilon R} s_{i\alpha_i} \leq e^{-\epsilon\Xi(1-\gamma/2)R} s_{i\alpha_i}$. Consequently, whenever $R \geq \bar{R}_0 := (\epsilon\Xi(1 - \gamma/2))^{-1} \log(\frac{TM}{\Gamma\nu})$,

$$q_i^{\tilde{\mathbf{p}}^{\epsilon-rlx}} \leq \frac{\sum_{t \neq \alpha_i} s_{it}}{s_{i\alpha_i}} \leq T e^{-\epsilon\Xi(1-\gamma/2)R} \leq \frac{\Gamma\nu}{M}.$$

This implies that, on $i \in \bar{\mathcal{S}}$

$$Y_i \cdot \mathbf{r}_i^\top \mathbf{v}^{mm} \geq 1 + \nu + Y_i \cdot (\mathbf{r}_i - \mathbf{x}_{i\alpha_i})^\top \mathbf{v}^{mm} \geq 1 + \nu - q_i M \|\mathbf{v}^{mm}\| \geq 1. \quad (80)$$

In words: Above a fixed $\bar{R}_0$ that only depends on $\gamma = \gamma(\delta)$, features $\mathbf{r}_i$ induced by all non-support indices $i \in \bar{\mathcal{S}}$ achieve margin at least 1. What remains is analyzing the margin shrinkage over the support vectors as in Theorem 5.

Controlling support margin and combining bounds: Over $\mathcal{S}$, suppose $\mathbf{v}^{mm}$ satisfies the SVM constraints on $\mathbf{r}_i$ with $Y_i \cdot \mathbf{r}_i^\top \mathbf{v}^{mm} \geq 1 - \epsilon_i/\Gamma$. Consequently, setting $\epsilon_{\max} = \sup_{i \in [n]} \epsilon_i$, $\mathbf{v}^{mm}$ achieves a label-margin of $\Gamma - \epsilon_{\max}$ on the dataset $(Y_i, \mathbf{r}_i)_{i \in [n]}$. Next, we recall the fact (79b) that worst-case perturbation is $\epsilon_{\max} \leq M \exp(-\Xi(1 - 0.5\gamma)R + \log T) = MT \exp(-\Xi(1 - 0.5\gamma)R)$. With this and (80), we upper bound the logistic loss of $(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{\epsilon-rlx})$ as follows.

$$\begin{aligned} \mathcal{L}(\tilde{\mathbf{v}}^{mm}, \tilde{\mathbf{p}}^{mm}) &\leq \max_{i \in [n]} \log(1 + \exp(-Y_i \mathbf{r}_i^\top \tilde{\mathbf{v}}^{mm})) \\ &\leq \max_{i \in [n]} \exp(-Y_i \mathbf{r}_i^\top \tilde{\mathbf{v}}^{mm}) \\ &\leq \exp(-r\Gamma + r\epsilon_{\max}) \\ &\leq e^{rMT \exp(-\Xi(1-0.5\gamma)R)} e^{-r\Gamma}. \end{aligned} \quad (81)$$

Conversely, we obtain a lower bound for $(\mathbf{v}_r, \mathbf{h}_R)$. Set $\mathbf{r}_i = \mathbf{X}_i^\top \mathbb{S}(\mathbf{K}_i \mathbf{h}_R)$. Recall the lower bound (79a) over the support vector set $\mathcal{S}$. Combining this with our Assumption C over the support vectors

of (SVM) implies that, solving (SVM) on $(Y_i, r_i)_{i \in [n]}$ achieves at most $\Gamma - \nu e^{-(1-\gamma)\Xi R}/T$ margin. Consequently, we have

$$\begin{aligned}
\mathcal{L}(\mathbf{v}_r, \mathbf{h}_R) &\geq \frac{1}{n} \max_{i \in [n]} \log(1 + \exp(-Y_i \mathbf{r}_i^\top \mathbf{v}_r)) \\
&\geq \frac{1}{2n} \max_{i \in [n]} \exp(-Y_i \mathbf{r}_i^\top \mathbf{v}_r) \wedge \log 2 \\
&\geq \frac{1}{2n} \exp(-r(\Gamma - \nu e^{-(1-\gamma)\Xi R}/T)) \wedge \log 2 \\
&\geq \frac{1}{2n} e^{r(v/T) \exp(-(1-\gamma)\Xi R)} e^{-r\Gamma} \wedge \log 2.
\end{aligned} \tag{82}$$

Observe that, this lower bound dominates the previous upper bound when R is large, namely, when (ignoring the multiplier $1/2n$ for brevity)

$$(v/T) e^{-(1-\gamma)\Xi R} \geq M T e^{-\Xi(1-0.5\gamma)R} \iff R \geq R_0 := \frac{2}{\gamma\Xi} \log\left(\frac{MT^2}{\nu}\right).$$

Thus, we obtain the desired contradiction since $\tilde{\mathbf{p}}^{\epsilon-r\Gamma}$ is a strictly better solution compared to $\mathbf{p}_R = \mathbf{h}_R$ (once R is sufficiently large).

• **Case 2: $\mathbf{v}_r/r$ does not converge.** This is the simpler scenario: There exists $\delta > 0$ such that we can find arbitrarily large r obeying $\|\mathbf{v}_r/r - \mathbf{v}^{mm}/\|\mathbf{v}^{mm}\|\| \geq \delta$. First, note that, due to the strong convexity of (SVM), for some $\gamma := \gamma(\delta) > 0$, $\mathbf{v}_r$ achieves a margin of at most $(\Gamma - \gamma)r$ on the dataset $(Y_i, \mathbf{x}_{i1})_{i \in [n]}$. By theorem's condition, we are provided that $\mathbb{S}(\mathbf{K}_i \mathbf{p}_R)_{\alpha_i} \rightarrow 1$. This immediately implies that, for any choice of $\epsilon = \gamma/3 > 0$, above some sufficiently large (r_0, R_0) , we have that $\|\mathbf{x}_i^{\mathbf{p}_R} - \mathbf{r}_i\| \leq \epsilon$. Following (81), this implies that, choosing $\tilde{\mathbf{v}}^{mm} = r\mathbf{v}^{mm}/\|\mathbf{v}^{mm}\|$ achieves a logistic loss of at most $e^{r\gamma/3} e^{-r\Gamma}$. Again using $\|\mathbf{x}_i^{\mathbf{p}_R} - \mathbf{r}_i\| \leq \epsilon$, for sufficiently large (r, R) we have that

$$\begin{aligned}
\min_{i \in [n]} Y_i \mathbf{v}_r^\top \mathbf{r}_i &\leq \min_{i \in [n]} Y_i \mathbf{v}_r^\top \mathbf{x}_{i1} + \sup_{i \in [n]} |\mathbf{v}_r^\top (\mathbf{r}_i - \mathbf{x}_{i1})| \\
&\leq (\Gamma - \gamma)r + \epsilon r \\
&\leq (\Gamma - 2\gamma/3)r.
\end{aligned}$$

This in turn implies that logistic loss is lower bounded by (following (82)),

$$\mathcal{L}(\mathbf{v}_r, \mathbf{p}_R) \geq \frac{1}{2n} e^{2\gamma r/3} e^{-r\Gamma} \wedge \log 2.$$

This dominates the above upper bound $e^{r\gamma/3} e^{-r\Gamma}$ of $\tilde{\mathbf{v}}^{mm}$ whenever $\frac{1}{2n} e^{2\gamma r/3} > 1 \iff r > \frac{3}{\gamma} \log(2n)$, (that is, when r is sufficiently large), again concluding the proof. ■

D Regularization Path of Attention with Nonlinear Head

So far our discussion has focused on the attention model with linear head. However, the conceptual ideas on optimal token selection via margin maximization also extends to a general nonlinear model under mild assumptions. The aim of this section is showcasing this generalization. Specifically, we consider the prediction model $f(\mathbf{X}) = \psi(\mathbf{X}^\top \mathbb{S}(\mathbf{K}\mathbf{p}))$ where $\psi(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}$ generalizes the linear head $\mathbf{v}$ of our attention model. For instance, following exposition in Section 1.1, $\psi(\cdot)$ can represent a multilayer transformer with $\mathbf{p}$ being a tunable prompt at the input layer. Recall that $(\mathbf{X}_i, \mathbf{K}_i, Y_i)_{i=1}^n$ is the dataset of the input-key-label tuples. We consider the training risk

$$\mathcal{L}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i, \psi(\mathbf{X}_i^\top \mathbf{s}_i^{\mathbf{p}})), \quad \text{where } \mathbf{s}_i^{\mathbf{p}} = \mathbb{S}(\mathbf{K}_i \mathbf{p}) \in \mathbb{R}^T. \tag{83}$$

The challenge with nonlinear $\psi(\cdot)$ is that, we lack a clear score function (Def. 1) unlike the previous sections. The assumption below introduces a generic condition that splits the tokens of each $\mathbf{X}_i$ into an *optimal* set O_i and *non-optimal* set $\bar{O}_i = [T] - O_i$. In words, non-optimal tokens are those that strictly increase the training risk $\mathcal{L}(\mathbf{p})$ if they are not fully suppressed by attention probabilities $\mathbf{s}_i^{\mathbf{p}}$.

Assumption D (Mixing non-optimal tokens hurt) *There exists sets $(O_i)_{i=1}^n \subset [T]$ as follows. Let $q_i^p = \sum_{t \in \bar{O}_i} s_{it}^p$ be the sum of softmax similarities over the non-optimal set for $\mathbf{p}$. Set $q_{\max}^p = \max_{i \in [n]} q_i^p$. For any $\Delta > 0$, there exists $\rho < 0$ such that:*

For all $\mathbf{p}, \mathbf{p}' \in \mathbb{R}^d$, if $\log(q_{\max}^p) \leq (1 + \Delta) \log(q_{\max}^{p'}) \wedge \rho$, then $\mathcal{L}(\mathbf{p}) < \mathcal{L}(\mathbf{p}')$.

This assumption is titled *mixing hurts* because the attention output $X_i^\top s_i^p$ is mixing the tokens of X_i and our condition is that, to achieve optimal risk, this mixture should not contain any non-optimal tokens. In particular, we require that, a model $\mathbf{p}$ that contains *exponentially less non-optimality* (quantified via $\log(q_{\max})$) compared to $\mathbf{p}'$ is strictly preferable. As we discuss in the supplementary material, Theorem 1 is in fact a concrete instance (with linear head $\mathbf{v}$) satisfying this condition.

Before stating our generic theorem, we need to introduce the max-margin separator towards which regularization path of attention will converge. This is a slightly general version of Section 2's (ATT-SVM) problem where we allow for a set of optimal tokens O_i for each input.

$$\mathbf{p}^{mm} = \arg \min_{\mathbf{p}} \|\mathbf{p}\| \quad \text{subject to} \quad \max_{\alpha \in O_i} \min_{\beta \in \bar{O}_i} \mathbf{p}^\top (\mathbf{k}_{i\alpha} - \mathbf{k}_{i\beta}) \geq 1, \quad \text{for all } i \in [n]. \quad (\text{ATT-SVM}')$$

Unlike (ATT-SVM), this problem is not necessarily convex when the optimal set O_i is not a singleton. To see this, imagine $n = d = 1$ and $T = 3$: Set the two optimal tokens as $\mathbf{k}_1 = 1$ and $\mathbf{k}_2 = -1$ and the non-optimal token as $\mathbf{k}_3 = 0$. The solution set of (ATT-SVM') is $\mathbf{p}^{mm} \in \{-1, 1\}$ whereas their convex combination $\mathbf{p} = 0$ violates the constraints. To proceed, our final result establishes the convergence of regularization path to the solution set of (ATT-SVM') under Assumption D.

Theorem 8 *Let $\mathcal{G}^{mm}$ be the set of global minima of (ATT-SVM'). Suppose its margin $\Xi := 1/\|\mathbf{p}^{mm}\| > 0$ and Assumption D holds. Let $\text{dist}(\cdot, \cdot)$ denote the ℓ_2 -distance between a vector and a set. Following (83), define $\bar{\mathbf{p}}(R) = \arg \min_{\|\mathbf{p}\| \leq R} \mathcal{L}(\mathbf{p})$. We have that $\lim_{R \rightarrow \infty} \text{dist}\left(\frac{\bar{\mathbf{p}}(R)}{\Xi R}, \mathcal{G}^{mm}\right) = 0$.*

We note that Theorem 1 is a corollary of this result where O_i 's and $\mathcal{G}^{mm}$ are singleton. Based on this result, with multiple optimal tokens, Theorem 1 would gracefully generalize to solve (ATT-SVM').

D.1 Proof of Theorem 8

Proof. The key idea is showing that, thanks to the exponential tail of softmax-attention, (harmful) contribution of the non-optimal token with the minimum margin can dominate the contribution of all other tokens as $R \rightarrow \infty$. This high-level approach is similar to earlier works on implicit bias of gradient descent with logistic loss [31, 22].

Pick $\mathbf{p}^{mm} \in \mathcal{G}^{mm}$ and set $\mathbf{p}_R^* = R \frac{\mathbf{p}^{mm}}{\|\mathbf{p}^{mm}\|}$. This will be the baseline model that $\mathbf{p}_R$ has to compete against. Also let $\bar{\mathbf{p}}_R = \frac{\mathbf{p}_R}{\Xi R}$. Now suppose $\text{dist}(\bar{\mathbf{p}}_R, \mathcal{G}^{mm}) \rightarrow 0$ as $R \rightarrow \infty$. Then, there exists $\delta > 0$ such that, we can always find arbitrarily large R obeying $\text{dist}(\bar{\mathbf{p}}_R, \mathcal{G}^{mm}) \geq \delta$.

Since $\bar{\mathbf{p}}_R$ is $\delta > 0$ bounded away from $\mathcal{G}^{mm}$, and $\|\bar{\mathbf{p}}_R\| = \|\mathbf{p}^{mm}\|$, $\bar{\mathbf{p}}_R$ strictly violates at least one of the inequality constraints in (ATT-SVM'). Otherwise, we would have $\bar{\mathbf{p}}_R \in \mathcal{G}^{mm}$. Without losing generality, suppose $\bar{\mathbf{p}}_R$ violates the first margin constraint, that is, for some $\gamma := \gamma(\delta) > 0$, $\max_{\alpha \in O_1} \min_{\beta \in \bar{O}_1} \bar{\mathbf{p}}_R^\top (\mathbf{k}_{1\alpha} - \mathbf{k}_{1\beta}) \leq 1 - \gamma$. Now, we will argue that this will lead to a contradiction as $R \rightarrow \infty$ since we will show that $\mathcal{L}(\mathbf{p}_R^*) < \mathcal{L}(\mathbf{p}_R)$ for sufficiently large R .

First, let us control $\mathcal{L}(\mathbf{p}_R^*)$. We study $s_i^* = \mathbb{S}(\mathbf{K}_i \mathbf{p}_R^*)$ and let $\alpha_i \in O_i$ be the index α in (ATT-SVM') for which $\text{margin}_i = \max_{\alpha \in O_i} \min_{\beta \in \bar{O}_i} (\mathbf{k}_{i\alpha} - \mathbf{k}_{i\beta})^\top \mathbf{p}^{mm} \geq 1$ is attained. Then, we bound the non-optimality amount q_i^* of $\mathbf{p}_R^*$ as

$$q_i^* = \frac{\sum_{t \in \bar{O}_i} \exp(\mathbf{k}_{it}^\top \mathbf{p}_R^*)}{\sum_{t \in [T]} \exp(\mathbf{k}_{it}^\top \mathbf{p}_R^*)} \leq \frac{\sum_{t \in \bar{O}_i} \exp(\mathbf{k}_{it}^\top \mathbf{p}_R^*)}{\exp(\mathbf{k}_{i\alpha_i}^\top \mathbf{p}_R^*)} \leq T \exp(-\Xi R).$$

Thus, $q_{\max}^* = \max_{i \in [n]} q_i^* \leq T \exp(-\Xi R)$. Secondly, we wish to control $\mathcal{L}(\mathbf{p}_R)$ by lower bounding the non-optimality in $\mathbf{p}_R$. Focusing on the first margin constraint, let $\alpha \in O_1$ be the index in (ATT-SVM') for which $\text{margin}_1 \leq 1 - \gamma$ is attained. Denoting the amount of non-optimality of the first input as $\hat{q}_1$, we find⁷

$$\hat{q}_1 = \frac{\sum_{t \in \bar{O}_1} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)}{\sum_{t \in [T]} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)} \geq \frac{1}{T} \frac{\sum_{t \in \bar{O}_1} \exp(\mathbf{k}_{1t}^\top \mathbf{p}_R)}{\exp(\mathbf{k}_{1\alpha}^\top \mathbf{p}_R)} \geq T^{-1} \exp(-\Xi R(1 - \gamma)).$$

⁷Here, we assumed margin is non-negative i.e. $\mathbf{k}_{1\alpha}^\top \mathbf{p}_R \geq \sup_{t \in \bar{O}_1} \mathbf{k}_{1t}^\top \mathbf{p}_R$. Otherwise, $\sup_{t \in [T]} \mathbf{k}_{1t}^\top \mathbf{p}_R$ is attained in $\bar{O}_1$ which implies $\hat{q}_1 \geq T^{-1}$. Thus, we can still use the identical inequality (84) with the choice $\gamma = 1$.

We similarly have $q_{\max}^* \geq T^{-1} \exp(-\Xi R)$. In conclusion, for $\mathbf{p}_R, \mathbf{p}_R^*$, denoting maximum non-optimality by $\hat{q}_{\max} \geq \hat{q}_1$ and $q_{\max}^*$, we respectively obtained

$$\begin{aligned} \log(\hat{q}_{\max}) &\geq -(1 - \gamma)(\Xi R) - \log T, \\ -(\Xi R) - \log T &\leq \log(q_{\max}^*) \leq -(\Xi R) + \log T. \end{aligned} \quad (84)$$

The above inequalities satisfy Assumption **D** as follows where $\mathbf{p} \leftarrow \mathbf{p}_R^*$ and $\mathbf{p}' \leftarrow \mathbf{p}_R$: Set $R_0 = 3\gamma^{-1}\Xi^{-1} \log T$ so that $\log T = \frac{\gamma\Xi R_0}{3}$. Secondly, set $\rho_0 = -\Xi R_0 - \log T$. This way, $\rho_0 \geq \log(q_{\max}^*)$ implies $R \geq R_0$ and $\log T \leq \frac{\gamma\Xi R}{3}$. Using the latter inequality, we bound the $\log T$ terms to obtain

- $\log(\hat{q}_{\max}) \geq -(1 - 2\gamma/3)(\Xi R)$, and
- $\log(q_{\max}^*) \leq -(1 - \gamma/3)(\Xi R)$.

To proceed, we pick $1 + \Delta = \frac{1-\gamma/3}{1-2\gamma/3}$ implying $\Delta := \frac{\gamma}{3-2\gamma}$. Finally, for this Δ , there exists $\rho(\Delta)$ which we need to ensure $\log(\hat{q}_{\max}) \leq \rho(\Delta)$. This can be guaranteed by picking sufficiently large R that ensures $\log(q_{\max}^*) \leq -(1 - \gamma/3)(\Xi R) \leq \rho(\Delta)$ to satisfy all conditions of Assumption **D**. Since such large R exists by initial assumption $\text{dist}(\bar{\mathbf{p}}_R, \mathcal{G}^{mm}) \rightarrow 0$, Assumption **D** in turn implies that $\mathcal{L}(\mathbf{p}_R^*) < \mathcal{L}(\mathbf{p}_R)$ contradicting with the optimality of $\mathbf{p}_R$ in (83). ■

D.2 Application to Linearly-mixed Labels

The following example shows that if non-optimal tokens result in reduced score (in terms of the alignment of prediction and label), Assumption **D** holds. The high-level idea behind this lemma is that, if the optimal risk is achieved by setting $q_{\max}^p = 0$, then, Assumption **D** will hold.

Lemma 14 (Linear label mixing) Recall $q_i^p = \sum_{t \in \bar{\mathcal{O}}_i} s_{it}^p$ from Assumption **D**. Suppose $Y_i \in \{-1, 1\}$ and

$$Y_i \cdot \psi(X_i^\top s_i^p) = v_i(1 - q_i^p) + Z_i,$$

for some $(v_i)_{i=1}^n > 0$. Here $Z_i = Z_i(\mathbf{p})$ is the contribution of non-optimal tokens to prediction. For some $C, \epsilon > 0$ and for all $\mathbf{p} \in \mathbb{R}^d$, assume

$$-Cq_{\max}^p \leq Z_i \leq (1 - \epsilon)v_i q_i^p. \quad (85)$$

Then, Assumption **D** holds for $\mathcal{L}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i \cdot \psi(X_i^\top s_i^p))$ when $\ell(\cdot)$ is a strictly decreasing loss function with continuous derivative.

Proof. Recall the assumption $Y_i \cdot \psi(X_i^\top s_i^p) = v_i(1 - q_i^p) + Z_i$ with Z_i obeying (85). Let us also write the loss function

$$\mathcal{L}(\mathbf{p}) = \frac{1}{n} \sum_{i=1}^n \ell(v_i(1 - s_i^p) + Z_i).$$

Define $q_{\max}^p = \sup_{i \leq [n]} q_i^p$. Let M be the maximum absolute value of score over tokens. Let

$$B = \max_{|x| \leq M} -\ell'(x) \geq A = \min_{|x| \leq M} -\ell'(x) > 0.$$

Through Taylor's Theorem (integral remainder), we have that

$$B(q_i^p v_i - Z_i) \geq \ell(v_i(1 - q_i^p) + Z_i) - \ell(v_i) \geq A(q_i^p v_i - Z_i) \geq \epsilon A v_i q_i^p.$$

Set $\mathcal{L}_\star = \frac{1}{n} \sum_{i=1}^n \ell(v_i)$. Set $C_+ = B(C + \max_{i \in [n]} v_i)$ and $C_- = n^{-1} A \epsilon \min_{i \in [n]} v_i$. This also implies

$$\begin{aligned} C_+ q_{\max}^p &\geq \frac{1}{n} \sum_{i \in [n]} B(q_i^p v_i - Z_i) \geq \mathcal{L}(\mathbf{p}) - \mathcal{L}_\star \geq \frac{1}{n} \sum_{i \in [n]} A(q_i^p v_i - Z_i) \\ &\geq \frac{1}{n} \sum_{i \in [n]} \epsilon A v_i q_i^p \geq C_- q_{\max}^p. \end{aligned}$$

Thus, to prove $\mathcal{L}(\mathbf{p}') > \mathcal{L}(\mathbf{p})$, we simply need to establish the stronger statement $C_- q_{\max}^{p'} > C_+ q_{\max}^p$.

Going back to the condition of Assumption D, any $\log(q_{\max}^p) \leq (1 + \Delta) \log(q_{\max}^{p'})$ obeys $q_{\max}^p \leq (q_{\max}^{p'})^{1+\Delta}$ i.e. $q_{\max}^{p'} \geq (q_{\max}^p)^{(1+\Delta)^{-1}}$. Following above, we wish to ensure $q_{\max}^{p'} > \Theta q_{\max}^p$ for such (p, p') pairs where $\Theta = C_+/C_- > 1$. This is guaranteed by

$$(q_{\max}^p)^{(1+\Delta)^{-1}-1} > \Theta \iff \frac{\Delta}{1+\Delta} \log(q_{\max}^p) < -\log(\Theta).$$

The above is satisfied by choosing a $\rho(\Delta) := -2(1 + \Delta^{-1}) \log(\Theta)$ in Assumption D. Thus, all p, p' with $\log(q_{\max}^p) \leq \rho = \rho(\Delta)$ satisfies the condition of Assumption D finishing the proof. ■

E Implementation Details and Additional Experiments

In this section, we provide implementation details and additional experiments.

- We build one attention layer using PyTorch. During training, we use SGD optimizer with learning rate 0.1 and train the model for 1000 iterations. To better visualize the convergence path, we normalize the gradient of p (and v) at each iteration.
- Next, given the gradient solution p , we determine locally-optimal indices to be those with the highest softmax scores. Using these optimal indices, we utilize python package `cvxopt` to build and solve (ATT-SVM), and then get solution p^{mm} . After obtaining p^{mm} , we also verify that these indices satisfy our local-optimal definition. The examples we use in the paper are all trivial to verify (by construction).
- In Figures 3(a) and 3(b), v^{mm} (blue dashed) is solved using python package `sklearn.svm` via (SVM) based on the given label information, and red dashed line represents p^{relax} direction instead, which is the solution of (10). Note that in both figures, $v^{mm}/\|v^{mm}\| = [0, 1]$. Therefore, in Figure 3(a) all optimal tokens are support vectors and $p^{relax} = p^{mm}$. Whereas in Figure 3(b), yellow ★ is not a support vector and only needs to satisfy positive correlation with p . Gray dashed line displays the p^{mm} direction.

Failure of gradient descent’s global convergence when $n \geq 2$ (refer to Theorem 2). Figure 8 provides a counter-example demonstrating that the $n = 1$ restriction is indeed necessary and tight to guarantee global convergence of the gradient descent iterates $p(t+1) = p(t) - \eta \nabla \mathcal{L}(p(t))$ on (ERM). For this example, we use logistic loss in (ERM). We set $n = T = d = 2$, implying that there is only one non-optimal token, thus Assumption B is satisfied. The red and blue lines represent GMM and LMM solutions, respectively. We note that the green star and teal square indicate the locally-optimal tokens. Specifically, referring to the local optimality definition (Definition 2), for LMM solution (p^{mm}) represented by the blue line, the square teal token does not have any SVM-neighbors. The arrows indicate the two trajectories originating from different initializations. This demonstrates that the gradient descent iterates $p(t+1) = p(t) - \eta \nabla \mathcal{L}(p(t))$ on (ERM) with two different initializations converge to two different SVM solutions (GMM and LMM). Results validate the necessity of $n = 1$ in Theorem 2 to provide the gradient descent convergence to p^{mm*} from any initialization.

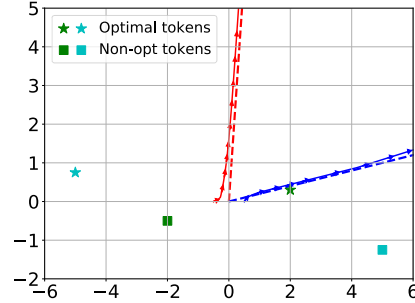


Figure 8: The convergence behavior of the gradient descent on the attention weights p using the logistic loss in (ERM) with $n = T = d = 2$.

The convergence behavior of gradient descent under over-parameterization. To illustrate Theorems 3 & 4, we have investigated the convergence behavior of $p(t)$ generated by gradient descent in Figure 9(a), using $n = 4$, $T = 6$, and conducted 1,000 random trials for varying $d \in \{2, 5, 10, 100, 300, 500\}$. These experiments use normalized gradient descent with learning rate 1 for 1000 iterations. Inputs x_{it} and the linear head v are uniformly sampled from the unit sphere, while Y_i is uniformly ± 1 , and W is set to I .

The bar plot in Figure 9(a) distinguishes between *non-saturated* softmax (red bars) and *saturated* softmax (other bars). Saturation is defined as average softmax probability over tokens selected by gradient descent are at least $1 - 10^{-5}$ and implies that attention selects one token per input. Note that, whenever the norm of gradient descent solution is finite, softmax will be non-saturated. For

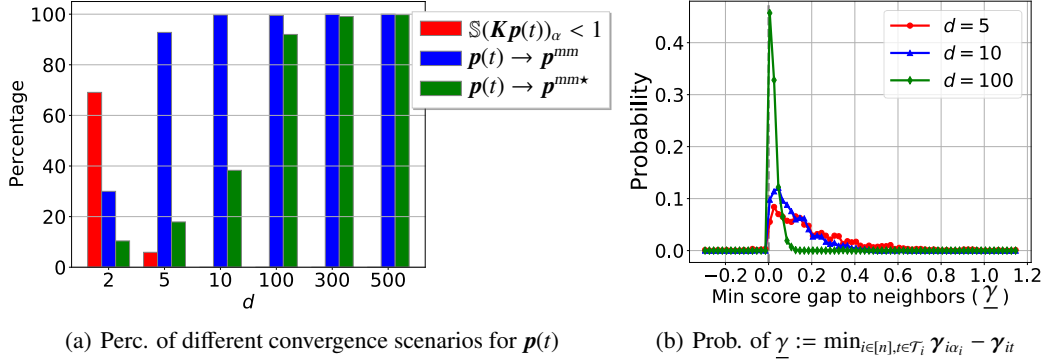


Figure 9: Convergence analysis of $p(t)$ trained with random data using gradient descent. (a) shows three scenarios: (1) attention failing to select one token per input (i.e. softmax is not saturated); (2) p converging towards p^{mm} ; and (3) p^{mm} equating to p^{mm*} with red, blue, and green bars, respectively. Considering saturated softmax instances where $p(t)$ selects one token α_i per-input, (b) presents histogram of the minimal score gap between α_i and its corresponding neighbors $\mathcal{T}_i$.

small d (e.g., $d = 2$), problem has small degrees of freedom to separate optimal tokens from the rest (i.e. no SVM solution for LMM directions) – especially due to label randomness. This results in a tall red bar capturing the finite-norm solutions. However, for larger d , we observe that softmax saturates (i.e. $\|p(t)\| \rightarrow \infty$) and we observe that the selected tokens α almost always converges to an LMM direction (blue bar) – this is in line with Theorems 3 & 4. We also study the convergence to the globally-optimal GMM which is represented by the green bar: GMM is a strict subset of LMM however as d increases, we observe that the probability of GMM convergence increases as well. This behavior is in line with what one would expect from over-parameterized deep learning theory [57, 58, 59, 60] and motivates future research. The average correlation coefficient between $p(t)$ and its associated LMM/GMM direction is 0.997, suggesting that, whenever softmax saturates, gradient descent indeed directionally converges to a LMM solution $p \in \mathcal{P}^{mm}$, confirming Theorem 4.

Furthermore, we found that there exist problem instances, with saturated softmax and $\|p(t)\| \rightarrow \infty$, that do not converge to either LMM or GMM. We analyzed this phenomenon using the minimum score gap, $\underline{\gamma} := \min_{i \in [n], t \in \mathcal{T}_i} \gamma_{i\alpha_i} - \gamma_{it}$, where $\mathcal{T}_i, i \in [n]$, represents the sets of SVM-neighbor tokens. Figure 9(b) provides the probability distribution of $\underline{\gamma}$ (with bins of width < 0.01) and demonstrates the rarity of such cases. Specifically, we found this happens less than 1% of the problems, that is, $\text{Prob}(\underline{\gamma} < 0) < 0.01$. Figure 9(b) also reveals that, in these scenarios, even if $\underline{\gamma} < 0$, it is typically close to zero i.e. even if there exists a SVM-neighbor with a higher score, it is only slightly so. This is not surprising since when token scores are close, we need a large number of gradient iterations to distinguish them. For all practical purposes, the optimization will treat both tokens equally and rather than solving (ATT-SVM), the more refined formulation (ATT-SVM') developed in Section D will be a better proxy. Confirming this intuition, we have verified that, over the instances $\underline{\gamma} < 0$, gradient descent solution is still > 0.99 correlated with the max-margin solution in average.

- In Figure 9, we again applied normalized gradient descent with a learning rate equal to 1. Each trial involved randomly generated data and training for 1000 iterations as discussed in Theorem 4. The tokens selected by p were denoted as $(\alpha_i)_{i=1}^n$, where $\alpha_i = \arg \max_{t \in [T]} \mathbb{S}(X_i p)_t$. The averaged softmax probabilities were calculated as $\bar{s} := \frac{1}{n} \sum_{i=1}^n \mathbb{S}(X_i p)_{\alpha_i}$ (same as Figure 3(c)). The red bars in Figure 9(b) represent the values of $\mathbb{P}(\bar{s} \leq 1 - 10^{-5})$ for each choice of d . Figure 10 displays the cumulative probability distribution of $\underline{\gamma}$ from Figure 9(b), with the gray dashed line indicating $\underline{\gamma} = 0$. From this figure, we observe that the minimal score gap exhibit a sharp transition at zero ($< 1\%$ of the instances have $\underline{\gamma} < 0$), demonstrating that, in most random problem instances with $\|p\| \rightarrow \infty$ ($\bar{s} \rightarrow 1$), problem directionally

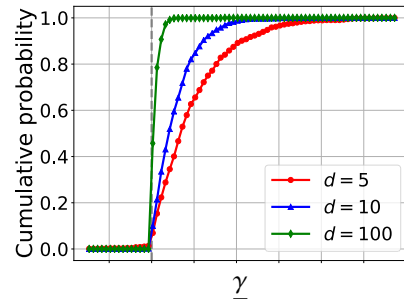


Figure 10: Cumulative prob. of the gap $\underline{\gamma} := \min_{i \in [n], t \in \mathcal{T}_i} \gamma_{i\alpha_i} - \gamma_{it}$ in Figure 9(b).

converges to an LMM i.e. $\mathbf{p}(t)/\|\mathbf{p}(t)\| \rightarrow \mathbf{p}^{mm}/\|\mathbf{p}^{mm}\|$. We believe the rare occurrence of a negative score gap is due to small score differences (so that optimal token is not clearly distinguished) and finite number of gradient iterations we run. Interestingly, even in the negative score gap scenarios, gradient descent is aligned with $\mathbf{p}^{mm}(\alpha)$ (even if $\mathbf{p}^{mm}(\alpha)$ is not LMM) which can be predicted from our Section D which handles the scenario where there are multiple optimal tokens per input.

F Addendum to Section 5

We provide an overview of the current literature on implicit regularization and attention mechanism.

F.1 Related Work on Implicit Regularization

The introduction of Support Vector Machines (SVM), which utilize explicit regularization to choose maximum margin classifiers, represents one of the earliest relevant literature in this field [61]. The concept of maximizing the margin was later connected to generalization performance [62]. From a practical perspective, exponential losses with decaying regularization exhibit asymptotic behavior similar to SVMs, as demonstrated in [22]. While the analysis of the perceptron [63] originally introduced the concept of margins, the method itself does not possess an inherent bias as it terminates with zero classification error. However, establishing a meaningful lower bound for the attained margin is not possible. Initial empirical investigations highlighting the implicit bias of descent methods focused on ℓ_1 -regularization, revealing that coordinate descent, when combined with the exponential loss, exhibits an inherent inclination towards ℓ_1 -regularized solutions [64].

This work draws extensively from the literature on implicit bias and regularization, which has provided valuable techniques and inspiration. A common observation in these studies is the convergence to a specific optimal solution over the training set. This phenomenon has been observed in various approaches, including coordinate descent [65, 66], gradient descent [30, 67, 25, 68, 69, 22, 70], deep linear networks [71, 72], ReLU networks [73, 74, 29, 75, 76], mirror descent [77, 78, 33, 36], and many others. The implicit bias of gradient descent in classification tasks involving separable data has been extensively examined by [22, 25, 26, 27, 28, 29]. The works on classification typically utilize logistic loss or exponentially-tailed losses to establish connections to margin maximization. The results have also been extended to non-separable data by [30, 31, 21]. Additionally, several papers have explored the implicit bias of stochastic gradient descent [37, 38, 41, 42], as well as adaptive and momentum-based methods [43, 44, 45, 46].

While there are some similarities between our optimization approach for $\mathbf{v}$ and existing works, the optimization of $\mathbf{p}$ presents notable differences. Firstly, our optimization problem is nonconvex and involves a composition of loss and softmax, which introduces new challenges and complexities. The presence of softmax adds a nonlinearity to the problem, requiring specialized techniques for analysis and optimization. Secondly, our analysis introduces the concept of locally-optimal tokens, which refers to tokens that achieve locally optimal solutions in their respective attention cones. This concept is crucial for understanding the behavior of the attention mechanism and its convergence properties. By focusing on the cones surrounding locally-optimal tokens, we provide a tailored analysis that captures the unique characteristics of the attention model. Overall, our work offers novel insights into the optimization of attention-based models and sheds light on the behavior of the attention mechanism during training.

F.2 Related Work on Attention Mechanism

As the backbone of Transformers [6], the self-attention mechanism [47, 48, 49, 50] plays a crucial role in computing feature representations by globally modeling long-range interactions within the input. Transformers have achieved remarkable empirical success in various domains, including natural language processing [4, 2], recommendation systems [79, 80, 81], and reinforcement learning [82, 83, 84]. With the introduction of Vision Transformer (ViT) [85], Transformer-based models [86, 87] have become a strong alternative to convolutional neural networks (CNN) and become prevalent in vision tasks.

However, the theoretical foundation of Transformers and self-attention mechanisms has remained largely unexplored. Some studies have established important results, including the Lipschitz constant of self-attention [88], properties of the neural tangent kernel [89, 90], and the expressive power and

Turing-completeness of Transformers [91, 92, 93, 51, 23, 94] with statistical guarantees [95, 96]. There is also a growing effort towards a theoretical understanding of emergent abilities of language models – such as in-context learning [97, 98, 99] and chain-of-thought [100, 101, 102] – which are inherently related to the models ability to attend to the relevant information within the input sequence.

Focusing on the self-attention component, Edelman et al. [51] theoretically shows that a single self-attention head can represent a sparse function of the input with a sample complexity for the generalization gap between the training loss and the test loss. However, they did not delve into the algorithmic aspects of training Transformers to achieve desirable loss. Sahiner et al. [52] and Ergen et al. [53] further explored the analysis of convex relaxations for self-attention, investigating potential optimization techniques and properties. The former work applies to self-attention with linear activation (rather than softmax) whereas the latter work attempts to approximate softmax via a linear operation with unit simplex constraints. In contrast, we directly study softmax and characterize its non-convex geometry. In terms of expressive ability, Baldi and Vershynin [54] investigated the capacity of attention layers to capture complex patterns and information, while Dong et al. [23] illustrates the propensity of attention networks to degenerate during the training process, with the result often being an output that is approximately a rank-1 matrix.

Recent works have made progress in characterizing the optimization and generalization dynamics of attention [55, 56, 103, 17, 104]. Jelassi et al. [55] studied gradient-based methods from random initialization and provided a theoretical analysis of the empirical finding that Vision Transformers learn position embeddings that recapitulate the spatial structure of the training data, even though this spatial structure is no longer explicitly represented after the image is split into patches. Li et al. [56] provided theoretical results on training three-layer ViTs for classification tasks. They quantified the importance of self-attention in terms of sample complexity for achieving zero generalization error, as well as the sparsity of attention maps when trained by stochastic gradient descent (SGD). In another related work, Nguyen et al. [104] proposed a primal-dual optimization framework that focuses on deriving attention as the dual expansion of a primal neural network layer. By solving a support vector regression problem, they gained a deeper understanding and explanation of various attention mechanisms. This framework also enables the creation of novel attention mechanisms, offering flexibility and customization in designing attention-based models. In another closely related work, Oymak et al. [17] analyzed the same attention model as ours, denoted by (ERM). Specifically, they jointly optimize $\mathbf{v}$, $\mathbf{p}$ for three gradient iterations for a contextual dataset model. This is in contrast to our emphasis on infinite-iteration behavior of $\mathbf{p}$ -only optimization. However, it is important to note that all of these works make certain assumptions about the data. Specifically, they assume that tokens are tightly clusterable or can be clearly split into relevant and irrelevant sets. Additionally, Li et al. [56] require specific assumptions on the initialization of the model, while Jelassi et al. [55] consider a simplified attention structure where the attention matrix is not directly parameterized with respect to the input.

In contrast, our work offers a comprehensive optimization-theoretic analysis of the attention model, establishing a formal connection to max-margin problems. While comparable works make assumptions on the dataset model, our results apply under minimal assumptions for general data and realistic conditions. Our analysis based on max-margin-equivalence allows us to gain a deeper understanding of the optimization geometry of attention and its behavior during the training process. As articulated in our experiments, our results lead to novel insights even for $n = 1, 2$ samples, $T = 2, 3$ tokens and $d = 2, 3$ dimensions (in contrast to [55, 56, 103, 17]). Notably, our work also presents the first theoretical understanding of the implicit bias exhibited by gradient descent methods in the context of the attention model. We remark that recent work [105] expands the theory presented in this work to 1-layer transformers. By uncovering the underlying optimization principles and thoroughly characterizing the directional convergence of attention, we provide valuable insights into the dynamics and generalization properties of attention-based models opening the path for future research.