

IMPLICIT REGULARIZATION FOR GROUP SPARSITY

Jiangyuan Li*, Thanh V. Nguyen, Chinmay Hegde[†] & Raymond K. W. Wong*

*Texas A&M University

[†]New York University

{jiangyuanli, raywong}@tamu.edu;
thanhng.cs@gmail.com; chinmay.h@nyu.edu

ABSTRACT

We study the implicit regularization of gradient descent towards structured sparsity via a novel neural reparameterization, which we call a “*diagonally grouped linear neural network*”. We show the following intriguing property of our reparameterization: gradient descent over the squared regression loss, without any explicit regularization, biases towards solutions with a group sparsity structure. In contrast to many existing works in understanding implicit regularization, we prove that our training trajectory cannot be simulated by mirror descent. We analyze the gradient dynamics of the corresponding regression problem in the general noise setting and obtain minimax-optimal error rates. Compared to existing bounds for implicit sparse regularization using diagonal linear networks, our analysis with the new reparameterization shows improved sample complexity. In the degenerate case of size-one groups, our approach gives rise to a new algorithm for sparse linear regression. Finally, we demonstrate the efficacy of our approach with several numerical experiments¹.

1 INTRODUCTION

Motivation. A salient feature of modern deep neural networks is that they are highly overparameterized with many more parameters than available training examples. Surprisingly, however, deep neural networks trained with gradient descent can generalize quite well in practice, even without explicit regularization. One hypothesis is that the dynamics of gradient descent-based training itself induce some form of implicit regularization, biasing toward solutions with low-complexity (Hardt et al., 2016; Neyshabur et al., 2017). Recent research in deep learning theory has validated the hypothesis of such implicit regularization effects. A large body of work, which we survey below, has considered certain (restricted) families of linear neural networks and established two types of implicit regularization — standard sparse regularization and ℓ_2 -norm regularization — depending on how gradient descent is initialized.

On the other hand, the role of *network architecture*, or the way the model is parameterized in implicit regularization, is less well-understood. Does there exist a parameterization that promotes implicit regularization of gradient descent towards richer structures beyond standard sparsity?

In this paper, we analyze a simple, prototypical hierarchical architecture for which gradient descent induces *group* sparse regularization. Our finding — that finer, *structured* biases can be induced via gradient dynamics — highlights the richness of co-designing neural networks along with optimization methods for producing more sophisticated regularization effects.

Background. Many recent theoretical efforts have revisited traditional, well-understood problems such as linear regression (Vaskevicius et al., 2019; Li et al., 2021; Zhao et al., 2019), matrix factorization (Gunasekar et al., 2018b; Li et al., 2018; Arora et al., 2019) and tensor decomposition (Ge et al., 2017; Wang et al., 2020), from the perspective of neural network training. For nonlinear models with squared error loss, Williams et al. (2019) and Jin & Montúfar (2020) study the implicit bias of gradient descent in wide depth-2 ReLU networks with input dimension 1. Other works (Gunasekar et al., 2018c; Soudry et al., 2018; Nacson et al., 2019) show that gradient descent biases the solution towards the max-margin (or minimum ℓ_2 -norm) solutions over separable data.

¹Code is available on <https://github.com/jiangyuan2li/Implicit-Group-Sparsity>

	NNs	Noise	Implicit vs. Explicit	Regularization
Vaskevicius et al. (2019)	DLNN	✓	Implicit (GD)	Sparsity
Dai et al. (2021)	LNN	✗	Explicit (ℓ_2 -penalty)	(Group) Quasi-norm
Jagadeesan et al. (2021)	LCNN	✗	Explicit (ℓ_2 -penalty)	Norm induced by SDP
Wu et al. (2020)	DLNN	✗	Implicit	ℓ_2 -norm
This paper	DGLNN	✓	Implicit (GD)	Structured sparsity

Table 1: Comparisons to related work on implicit and explicit regularization. Here, GD stands for gradient descent, (D)LNN/CNN for (diagonal) linear/convolutional neural network, and DGLNN for diagonally grouped linear neural network.

Outside of implicit regularization, several other works study the inductive bias of network architectures under *explicit* ℓ_2 regularization on model weights (Pilanci & Ergen, 2020; Sahiner et al., 2020). For multichannel linear convolutional networks, Jagadeesan et al. (2021) show that ℓ_2 -norm minimization of weights leads to a norm regularizer on predictors, where the norm is given by a semidefinite program (SDP). The representation cost in predictor space induced by explicit ℓ_2 regularization on (various different versions of) linear neural networks is studied in Dai et al. (2021), which demonstrates several interesting (induced) regularizers on the linear predictors such as ℓ_p quasi-norms and group quasi-norms. However, these results are silent on the behavior of gradient descent-based training *without* explicit regularization. In light of the above results, we ask the following question:

Beyond ℓ_2 -norm, sparsity and low-rankness, can gradient descent induce other forms of implicit regularization?

Our contributions. In this paper, we rigorously show that a *diagonally-grouped linear neural network* (see Figure 1b) trained by gradient descent with (proper/partial) weight normalization induces *group-sparse* regularization: a form of structured regularization that, to the best of our knowledge, has not been provably established in previous work.

One major approach to understanding implicit regularization of gradient descent is based on its equivalence to a mirror descent (on a different objective function) (e.g., Gunasekar et al., 2018a; Woodworth et al., 2020). However, we show that, for the diagonally-grouped linear network architecture, the gradient dynamics is beyond mirror descent. We then analyze the convergence of gradient flow with early stopping under orthogonal design with possibly noisy observations, and show that the obtained solution exhibits an implicit regularization effect towards structured (specifically, group) sparsity. In addition, we show that weight normalization can deal with instability related to the choices of learning rates and initialization. With weight normalization, we are able to obtain a similar implicit regularization result but in more general settings: orthogonal/non-orthogonal designs with possibly noisy observations. Also, the obtained solution can achieve minimax-optimal error rates.

Overall, compared to existing analysis of diagonal linear networks, our model design — that induces structured sparsity — exhibits provably improved sample complexity. In the degenerate case of size-one groups, our bounds coincide with previous results, and our approach can be interpreted as a new algorithm for sparse linear regression.

Our techniques. Our approach is built upon the *power reparameterization* trick, which has been shown to promote model sparsity (Schwarz et al., 2021). Raising the parameters of a linear model element-wisely to the N -th power ($N > 1$) results in that parameters of smaller magnitude receive smaller gradient updates, while parameters of larger magnitude receive larger updates. In essence, this leads to a “rich get richer” phenomenon in gradient-based training. In Gissin et al. (2019) and Berthier (2022), the authors analyze the gradient dynamics on a toy example, and call this “incremental learning”. Concretely, for a linear predictor $\mathbf{w} \in \mathbb{R}^p$, if we re-parameterize the model as $\mathbf{w} = \mathbf{u}^{\circ N} - \mathbf{v}^{\circ N}$ (where $\mathbf{u}^{\circ N}$ means the N -th element-wise power of $\mathbf{u}$), then gradient descent will bias the training towards sparse solutions. This reparameterization is equivalent to a diagonal linear network, as shown in Figure 1a. This is further studied in Woodworth et al. (2020) for interpolating predictors, where they show that a small enough initialization induces ℓ_1 -norm regularization. For noisy settings, Vaskevicius et al. (2019) and Li et al. (2021) show that gradient descent converges to sparse models with early stopping. In the special case of sparse recovery from under-sampled

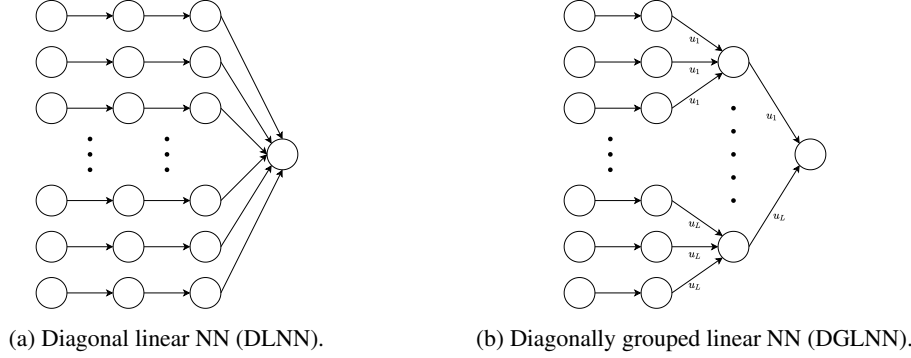


Figure 1: An illustration of the two architectures for standard and group sparse regularization.

observations (or compressive sensing), the optimal sample complexity can also be obtained via this reparameterization (Chou et al., 2021).

Inspired by this approach, we study a novel model reparameterization of the form $\mathbf{w} = [\mathbf{w}_1, \dots, \mathbf{w}_L]$, where $\mathbf{w}_l = u_l^2 \mathbf{v}_l$ for each group $l \in \{1, \dots, L\}$. (One way to interpret this model is to think of u_l as the “magnitude” and $\mathbf{v}_l$ as the “direction” of the subvector corresponding to each group; see Section 2 for details.) This corresponds to a special type of linear neural network architecture, as shown in Figure 1b. A related architecture has also been recently studied in Dai et al. (2021), but there the authors have focused on the bias induced by an *explicit* ℓ_2 regularization on the weights and have not investigated the effect of gradient dynamics.

The diagonally linear network parameterization of Woodworth et al. (2020); Li et al. (2021) does not suffer from identifiability issues. In contrast to that, in our setup the “magnitude” parameter u_l of each group interacts with the norm of the “direction”, $\|\mathbf{v}_l\|_2$, causing a fundamental problem of identifiability. By leveraging the layer balancing effect (Du et al., 2018) in DGLNN, we verify the group regularization effect implicit in gradient flow with early stopping. But gradient flow is idealized; for a more practical algorithm, we use a variant of gradient descent based on *weight normalization*, proposed in (Salimans & Kingma, 2016), and studied in more detail in (Wu et al., 2020). Weight normalization has been shown to be particularly helpful in stabilizing the effect of learning rates (Morwani & Ramaswamy, 2022; Van Laarhoven, 2017). With weight normalization, the learning effect is separated into magnitudes and directions. We derive the gradient dynamics on both magnitudes and directions with perturbations. Directions guide magnitude to grow, and as the magnitude grows, the directions get more accurate. Thereby, we are able to establish regularization effect implied by such gradient dynamics.

A remark on grouped architectures. Finally, we remark that grouping layers have been commonly used in grouped CNN and grouped attention mechanisms (Xie et al., 2017; Wu et al., 2021), which leads to parameter efficiency and better accuracy. Group sparsity is also useful for deep learning models in multi-omics data for survival prediction (Xie et al., 2019). We hope our analysis towards diagonally grouped linear NN could lead to more understanding of the inductive biases of grouping-style architectures.

2 SETUP

Notation. Denotes the set $\{1, 2, \dots, L\}$ by $[L]$, and the vector ℓ_2 norm by $\|\cdot\|$. We use $\mathbf{1}_p$ and $\mathbf{0}_p$ to denote p -dimensional vectors of all 1s and all 0s correspondingly. Also, $\odot$ represents the entry-wise multiplication whereas $\beta^{\odot N}$ denotes element-wise power N of a vector β . We use $\mathbf{e}_i$ to denote the i^{th} canonical vector. We write inequalities up to multiplicative constants using the notation $\lesssim$, whereby the constants do not depend on any problem parameter.

Observation model. Suppose that the index set $[p] = \cup_{l=1}^L G_l$ is partitioned into L disjoint (i.e., non-overlapping) groups $G_1, G_2, \dots, G_L$ where $G_i \cap G_j = \emptyset, \forall i \neq j$. The size of G_l is denoted by $p_l = |G_l|$ for $l \in [L]$. Let $\mathbf{w}^* \in \mathbb{R}^p$ be a p -dimensional vector where the entries of $\mathbf{w}^*$ are non-zero only on a subset of groups. We posit a linear model of data where observations $(\mathbf{x}_i, y_i) \in \mathbb{R}^p \times \mathbb{R}$, $i \in$

$[n]$ are given such that $y_i = \langle \mathbf{x}_i, \mathbf{w}^* \rangle + \xi_i$ for $i = 1, \dots, n$, and $\boldsymbol{\xi} = [\xi_1, \dots, \xi_n]^\top$ is a noise vector. Note that we do not impose any special restriction between n (the number of observations) and p (the dimension). We write the linear model in the following matrix-vector form: $\mathbf{y} = \mathbf{X}\mathbf{w}^* + \boldsymbol{\xi}$, with the $n \times p$ design matrix $\mathbf{X} = [\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_L]$, where $\mathbf{X}_l \in \mathbb{R}^{n \times p_l}$ represents the features from the l^{th} group G_l , for $l \in [L]$. We make the following assumptions on $\mathbf{X}$:

Assumption 1. *The design matrix $\mathbf{X}$ satisfies*

$$\sup_{\|\beta_1\| \leq 1, \|\beta_2\| \leq 1} \left| \left\langle \beta_1, \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) \beta_2 \right\rangle \right| \leq \delta_{in}, \quad \text{where } \beta_1, \beta_2 \in \mathbb{R}^{p_l}, \quad (1)$$

and

$$\sup_{\|\beta_1\| \leq 1, \|\beta_2\| \leq 1} \left| \left\langle \frac{1}{\sqrt{n}} \mathbf{X}_l \beta_1, \frac{1}{\sqrt{n}} \mathbf{X}_{l'} \beta_2 \right\rangle \right| \leq \delta_{out}, \quad \text{where } \beta_1 \in \mathbb{R}^{p_l}, \beta_2 \in \mathbb{R}^{p_{l'}}, l \neq l', \quad (2)$$

for some constants $\delta_{in}, \delta_{out} \in (0, 1)$.

The first part (1) is a within-group eigenvalue condition while the second part (2) is a between-group block coherence assumption. There are multiple ways to construct a sensing matrix to fulfill these two conditions (Eldar & Bolcskei, 2009; Baraniuk et al., 2010). One of them is based on the fact that random Gaussian matrices satisfy such conditions with high probability (Stojnic et al., 2009).

Reparameterization. Our goal is to learn a parameter $\mathbf{w}$ from the data $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ with coefficients which obey group structure. Instead of imposing an explicit group-sparsity constraint on $\mathbf{w}$ (e.g., via weight penalization by group), we show that gradient descent on the *unconstrained* regression loss can still learn $\mathbf{w}^*$, provided we design a special reparameterization. Define a mapping $g(\cdot) : [p] \rightarrow [L]$ from each index i to its group $g(i)$. Each parameter is rewritten as $w_i = u_{g(i)}^2 v_i, \forall i \in [p]$. The parameterization $G(\cdot) : \mathbb{R}_+^L \times \mathbb{R}^p \rightarrow \mathbb{R}^p$ reads

$$[u_1, \dots, u_L, v_1, v_2, \dots, v_p] \rightarrow [u_1^2 v_1, u_1^2 v_2, \dots, u_L^2 v_p].$$

This corresponds to the 2-layer neural network architecture displayed in Figure 1b, in which $\mathbf{W}_1 = \text{diag}(v_1, \dots, v_p)$, and $\mathbf{W}_2$ is “diagonally” tied within each group:

$$\mathbf{W}_2 = \text{diag}(u_1, \dots, u_1, u_2, \dots, u_2, \dots, u_L, \dots, u_L).$$

Gradient dynamics. We learn $\mathbf{u}$ and $\mathbf{v}$ by minimizing the standard squared loss:

$$\mathcal{L}(\mathbf{u}, \mathbf{v}) = \frac{1}{2} \|\mathbf{y} - \mathbf{X}[(\mathbf{D}\mathbf{u})^{\odot 2} \odot \mathbf{v}]\|^2,$$

where

$$\mathbf{D} = \begin{pmatrix} \mathbf{1}_{p_1} & \mathbf{0}_{p_1} & \dots & \mathbf{0}_{p_1} \\ \mathbf{0}_{p_2} & \mathbf{1}_{p_2} & \dots & \mathbf{0}_{p_2} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0}_{p_L} & \mathbf{0}_{p_L} & \dots & \mathbf{1}_{p_L} \end{pmatrix} \in \mathbb{R}^{p \times L}.$$

By simple algebra, the gradients with respect to $\mathbf{u}$ and $\mathbf{v}$ read as follows:

$$\begin{aligned} \nabla_{\mathbf{u}} L &= 2\mathbf{D}^\top (\mathbf{v} \odot [\mathbf{X}^\top \mathbf{X}((\mathbf{D}\mathbf{u})^{\odot 2} \odot \mathbf{v} - \mathbf{w}^*) - \mathbf{X}^\top \boldsymbol{\xi}] \odot \mathbf{D}\mathbf{u}), \\ \nabla_{\mathbf{v}} L &= [\mathbf{X}^\top \mathbf{X}((\mathbf{D}\mathbf{u})^{\odot 2} \odot \mathbf{v} - \mathbf{w}^*) - \mathbf{X}^\top \boldsymbol{\xi}] \odot (\mathbf{D}\mathbf{u})^{\odot 2}. \end{aligned}$$

Denote $\mathbf{r}(t) = \mathbf{y} - \sum_{l=1}^L u_l^2(t) \mathbf{X}_l \mathbf{v}_l(t)$. For each group $l \in [L]$, the gradient flow reads

$$\frac{\partial u_l(t)}{\partial t} = \frac{2}{n} u_l(t) \mathbf{v}_l^\top(t) \mathbf{X}_l^\top \mathbf{r}(t), \quad \frac{\partial \mathbf{v}_l(t)}{\partial t} = \frac{1}{n} u_l^2(t) \mathbf{X}_l^\top \mathbf{r}(t). \quad (3)$$

Although we are not able to transform the gradient dynamics back onto $\mathbf{w}(t)$ due to the overparameterization, the extra term $u_l(t)$ on group magnitude leads to “incremental learning” effect.

3 ANALYSIS OF GRADIENT FLOW

3.1 FIRST ATTEMPT: MIRROR FLOW

Existing results about implicit bias in overparameterized models are mostly based on recasting the training process from the parameter space $\{\mathbf{u}(t), \mathbf{v}(t)\}_{t \geq 0}$ to the predictor space $\{\mathbf{w}(t)\}_{t \geq 0}$ (Woodworth et al., 2020; Gunasekar et al., 2018a). If properly performed, the (induced) dynamics in the predictor space can now be analyzed by a classical algorithm: mirror descent (or mirror flow). Implicit regularization is demonstrated by showing that the limit point satisfies a KKT (Karush–Kuhn–Tucker) condition with respect to minimizing some regularizer $R(\cdot)$ among all possible solutions.

At first, we were unable to express the gradient dynamics in Eq. (3) in terms of $\mathbf{w}(t)$ (i.e., in the predictor space), due to complicated interactions between $\mathbf{u}$ and $\mathbf{v}$. This hints that the training trajectory induced by an overparameterized DGLNN may not be analyzed by mirror flow techniques. In fact, we prove a stronger negative result, and rigorously show that the corresponding dynamics *cannot* be recast as a mirror flow. Therefore, we conclude that our subsequent analysis techniques are necessary and do not follow as a corollary from existing approaches.

We first list two definitions from differential topology below.

Definition 1. Let M be a smooth submanifold of $\mathbb{R}^D$. Given two C^1 vector fields of X, Y on M , we define the Lie Bracket of X and Y as $[X, Y](x) := \partial Y(x)X(x) - \partial X(x)Y(x)$.

Definition 2. Let M be a smooth submanifold of $\mathbb{R}^D$. A C^2 parameterization $G : M \rightarrow \mathbb{R}^d$ is said to be commuting iff for any $i, j \in [d]$, the Lie Bracket $[\nabla G_i, \nabla G_j](x) = 0$ for all $x \in M$.

The parameterization studied in most existing works on diagonal networks is separable, meaning that each parameter only affects one coordinate in the predictor space. In DGLNN, the parameterization is not separable, due to the shared parameter $\mathbf{u}$ within each group. We formally show that it is indeed not commuting.

Lemma 1. $G(\cdot)$ is not a commuting parameterization.

Non-commutativity of the parameterization implies that moving along $-\nabla G_i$ and then $-\nabla G_j$ is different with moving with $-\nabla G_j$ first and then $-\nabla G_i$. This causes extra difficulty in analyzing the gradient dynamics. Li et al. (2022) study the equivalence between gradient flow on reparameterized models and mirror flow, and show that a commuting parameterization is a sufficient condition for when a gradient flow with certain parameterization simulates a mirror flow. A complementary necessary condition is also established on the Lie algebra generated by the gradients of coordinate functions of G with order higher than 2. We show that the parameterization $G(\cdot)$ violates this necessary condition.

Theorem 1. There exists an initialization $[\mathbf{u}_{init}^\top, \mathbf{v}_{init}^\top] \in \mathbb{R}_+^L \times \mathbb{R}^p$ and a time-dependent loss L_t such that gradient flow under $L_t \odot G$ starting from $[\mathbf{u}_{init}^\top, \mathbf{v}_{init}^\top]$ cannot be written as a mirror flow with respect to any Legendre function R under the loss L_t .

The detailed proof is deferred to the Appendix. Theorem 1 shows that the gradient dynamics implied in DGLNN cannot be emulated by mirror descent. Therefore, a different technique is needed to analyze the gradient dynamics and any associated implicit regularization effect.

3.2 LAYER BALANCING AND GRADIENT FLOW

Let us first introduce relevant quantities. Following our reparameterization, we rewrite the true parameters for each group l as

$$\mathbf{w}_l^* = (u_l^*)^2 \mathbf{v}_l^*, \quad \|\mathbf{v}_l^*\|_2 = 1, \quad \mathbf{v}_l^* \in \mathbb{R}^{p_l}.$$

The support is defined on the group level, where $S = \{l \in [L] : u_l^* > 0\}$ and the support size is defined as $s = |S|$. We denote $u_{max}^* = \max\{u_l^* | l \in S\}$, and $u_{min}^* = \min\{u_l^* | l \in S\}$.

The gradient dynamics in our reparameterization does not preserve $\|\mathbf{v}_l(t)\|_2 = 1$, which causes difficulty to identify the magnitude of each u_l and $\|\mathbf{v}_l(t)\|_2$. Du et al. (2018) and Arora et al. (2018)

show that the gradient flow of multi-layer homogeneous functions effectively enforces the differences between squared norms across different layers to remain invariant. Following the same idea, we discover a similar balancing effect in DGLNN between the parameter $\mathbf{u}$ and $\mathbf{v}$.

Lemma 2. For any $l \in [L]$, we have

$$\frac{d}{dt} \left(\frac{1}{2} u_l^2 - \|\mathbf{v}_l\|^2 \right) = 0.$$

The balancing result eliminates the identifiability issue on the magnitudes. As the coordinates within one group affect each other, the direction which controls the growth rate of both $\mathbf{u}$ and $\mathbf{v}$ need to be determined as well.

Lemma 3. If the initialization $\mathbf{v}_l(0)$ is proportional to $\frac{1}{n} \mathbf{X}_l^\top \mathbf{y}$, then

$$\left\langle \frac{\mathbf{v}_l(0)}{\|\mathbf{v}_l(0)\|}, \mathbf{v}_l^* \right\rangle \geq 1 - \left(\delta_{in} + L\delta_{out} + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_2 / (u_l^*)^2 \right)^2.$$

Note that this initialization can be obtained by a single step of gradient descent with $\mathbf{0}$ initialization. Lemma 3 suggests the direction is close to the truth at the initialization. We can further normalize it to be $\|\mathbf{v}_l(0)\|_2^2 = \frac{1}{2} u_l^2(0)$ based on the balancing criterion. The magnitude equality, $\|\mathbf{v}_l(t)\|_2^2 = \frac{1}{2} u_l^2(t)$, is preserved by Lemma 2. However, ensuring the closeness of the direction throughout the gradient flow presents significant technical difficulties. That said, we are able to present a meaningful implicit regularization result of the gradient flow under orthogonal (and noisy) settings.

Theorem 2. Fix $\epsilon > 0$. Consider the case where $\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l = \mathbf{I}$, $\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} = \mathbf{O}$, $l \neq l'$, the initialization $u_l(0) = \theta < \frac{\epsilon}{2(u_{max}^*)^2}$ and $\mathbf{v}_l(0) = \eta_l \frac{1}{n} \mathbf{X}_l^\top \mathbf{y}$ with $\|\mathbf{v}_l(0)\|_2^2 = \frac{1}{2} \theta^2$, $\forall l \in [L]$, there exists an lower bound and upper bound of the time $T_l < T_u$ in the gradient flow in Eq. (3), such that for any $T_l \leq t \leq T_u$ we have

$$\|u_l^2(t) \mathbf{v}_l(t) - \mathbf{w}_l^*\|_\infty \lesssim \begin{cases} \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty \vee \epsilon, & \text{if } l \in S. \\ \theta^{3/2}, & \text{if } l \notin S. \end{cases}$$

Theorem 2 states the error bounds for the estimation of the *true* weights $\mathbf{w}^*$. For entries outside the (true) support, the error is controlled by $\theta^{3/2}$. When θ is small, the algorithm keeps all non-supported entries to be close to zero through iterations while maintaining the guarantee for supported entries. Theorem 2 shows that under the assumption of orthogonal design, gradient flow with early stopping is able to obtain the solution with group sparsity.

4 GRADIENT DESCENT WITH WEIGHT NORMALIZATION

Algorithm 1 Gradient descent with weight normalization

Initialize: $\mathbf{u}(0) = \alpha \mathbf{1}$, unit norm initialization $\mathbf{v}_l(0)$ for each $l \in [L]$, $\eta_{l,t} = \frac{1}{u_l^4(t)}$.
for $t = 0$ to T **do**
 $\mathbf{z}(t+1) = \mathbf{v}(t) - \eta_{l,t} \nabla_{\mathbf{v}} \mathcal{L}(\mathbf{u}(t), \mathbf{v}(t))$
 $\mathbf{v}_l(t+1) = \frac{\mathbf{z}_l(t+1)}{\|\mathbf{z}_l(t+1)\|_2}, \forall l \in [L]$
 $\mathbf{u}(t+1) = \mathbf{u}(t) - \gamma \nabla_{\mathbf{u}} \mathcal{L}(\mathbf{u}(t), \mathbf{v}(t+1))$
 if the early stopping criterion is satisfied **then**
 stop
 end if
end for

We now seek a more practical algorithm with more general assumptions and requirements on initialization. To speed up the presentation, we will directly discuss the corresponding variant of (the more practical) gradient descent instead of gradient flow. When standard gradient descent is

applied on DGLNN, initialization for directions is very crucial; The algorithm may fail even with a very small initialization when the direction is not accurate, as shown in Appendix E. The balancing effect (Lemma 2) is sensitive to the step size, and errors may accumulate (Du et al., 2018).

Weight normalization as a commonly used training technique has been shown to be helpful in stabilizing the training process. The identifiability of the magnitude is naturally resolved by weight normalization on each $\mathbf{v}_l$. Moreover, weight normalization allows for a larger step size on $\mathbf{v}$, which makes the direction estimation at each step behave like that at the origin point. This removes the restrictive assumption of orthogonal design. With these intuitions in mind, we study the gradient descent algorithm with weight normalization on $\mathbf{v}$ summarized in Algorithm 1. One advantage of our algorithm is that it converges with *any* unit norm initialization $\mathbf{v}_l(0)$. The step size on $\mathbf{u}(t)$ is chosen to be small enough in order to enable the incremental learning, whereas the step size on $\mathbf{v}(t)$ is chosen as $\eta_{l,t} = \frac{1}{u_l^4(t)}$ as prescribed by our theoretical investigation. For convenience, we define

$$\zeta = 80 \left(\left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \vee \epsilon \right),$$

for a precision parameter $\epsilon > 0$. The convergence of Algorithm 1 is formalized as follows:

Theorem 3. Fix $\epsilon > 0$. Consider Algorithm 1 with

$$u_l(0) = \alpha < \frac{\epsilon^4 \wedge 1}{(u_{max}^*)^8} \wedge \frac{1}{80L} (u_{min}^*)^2 \wedge \frac{\epsilon}{L}, \quad \forall l \in [L],$$

any unit-norm initialization on $\mathbf{v}_l$ for each $l \in [L]$ and $\gamma \leq \frac{1}{20(u_{max}^*)^2}$. Suppose Assumption 1 is satisfied with $\delta_{in} \leq \frac{(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{(u_{min}^*)^2}{120s(u_{max}^*)^2}$. There exist a lower bound on the number of iterations

$$T_{lb} = \frac{\log \frac{(u_{max}^*)^2}{2\alpha^2}}{2 \log(1 + \frac{\gamma}{2}(\zeta \vee (u_{min}^*)^2))} + \left\lceil \log_2 \frac{(u_{max}^*)^2}{\zeta} \right\rceil \frac{5}{2\gamma(\zeta \vee (u_{min}^*)^2)},$$

and an upper bound

$$T_{ub} \geq \frac{5}{16\gamma(\zeta \vee (u_{min}^*)^2)} \log \frac{1}{\alpha^4},$$

such that $T_{lb} \leq T_{ub}$ and for any $T_{lb} \leq t \leq T_{ub}$,

$$\|u_l^2(t)\mathbf{v}_l(t) - \mathbf{w}_l^*\|_\infty \lesssim \begin{cases} \left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \vee \epsilon, & \text{if } l \in S \\ \alpha, & \text{if } l \notin S \end{cases}.$$

Similarly as Theorem 2, Theorem 3 states the error bounds for the estimation of the *true* weights $\mathbf{w}^*$. When α is small, the algorithm keeps all non-supported entries to be close to zero through iterations while maintaining the guarantee for supported entries. Compared to the works on implicit (unstructured) sparse regularization (Vaskevicius et al., 2019; Chou et al., 2021), our assumption on the incoherence parameter δ_{out} scales with $1/s$, where s is the number of non-zero groups, instead of the total number of non-zero entries. Therefore, the relaxed bound on δ_{out} implies an improved sample complexity, which is also observed experimentally in Figure 4. We now state a corollary in a common setting with independent random noise, where (asymptotic) recovery of $\mathbf{w}^*$ is possible.

Definition 3. A random variable Y is σ -sub-Gaussian if for all $t \in \mathbb{R}$ there exists $\sigma > 0$ such that

$$\mathbb{E} e^{tY} \leq e^{\sigma^2 t^2 / 2}.$$

Corollary 1. Suppose the noise vector $\boldsymbol{\xi}$ has independent σ^2 -sub-Gaussian entries and $\epsilon = 2\sqrt{\frac{\sigma^2 \log(2p)}{n}}$. Under the assumptions of Theorem 3, Algorithm 1 produces $\mathbf{w}(t) = (\mathbf{D}\mathbf{u}(t))^{\odot 2} \odot \mathbf{v}(t)$ that satisfies $\|\mathbf{w}(t) - \mathbf{w}^*\|_2^2 \lesssim (s\sigma^2 \log p)/n$ with probability at least $1 - 1/(8p^3)$ for any t such that $T_{lb} \leq t \leq T_{ub}$.

Note that the error bound we obtain is minimax-optimal. Despite these appealing properties of Algorithm 1, our theoretical results require a large step size on each $\mathbf{v}_l(t)$, which may cause instability at later stages of learning. We observe this instability numerically (see Figure 6, Appendix E). Although

the estimation error of $\mathbf{w}^*$ remains small (which aligns with our theoretical result), individual entries in $\mathbf{v}$ may fluctuate considerably. Indeed, the large step size is mainly introduced to maintain a strong directional information extracted from the gradient of $\mathbf{v}_l(t)$ so as to stabilize the updates of $\mathbf{u}(t)$ at the early iterations. Therefore, we also propose Algorithm 2, a variant of Algorithm 1, where we decrease the step size after a certain number of iterations.

Algorithm 2. Run Algorithm 1 with the same setup till each $u_l(t), l \in [L]$ gets roughly accurate, set $\eta_{l,t} = \eta$. Continue Algorithm 1 until early stopping criterion is satisfied.

Theorem 4. Under the assumptions of Theorem 3 with replacing the condition on δ 's by $\delta_{in} \leq \frac{\sqrt{\zeta}(u_{min}^*)^2}{120(u_{max}^*)^3}$ and $\delta_{out} \leq \frac{\sqrt{\zeta}(u_{min}^*)^2}{120s(u_{max}^*)^3}$, we apply Algorithm 2 with $\eta_{l,t} = \frac{1}{u^4(t)}$ at the beginning, and $\eta_{l,t} = \eta \leq \frac{4}{9(u_{max}^*)^2}$ after $\forall l \in [L], u_l^2(t) \geq \frac{1}{2}(u_l^*)^2$, then with the same T_{lb} and T_{ub} , we have that for any $T_{lb} \leq t \leq T_{ub}$,

$$\|u_l^2(t)\mathbf{v}_l(t) - \mathbf{w}_l^*\|_\infty \lesssim \begin{cases} \left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \vee \epsilon, & \text{if } l \in S. \\ \alpha, & \text{if } l \notin S. \end{cases}$$

In Theorem 4, the criterion to decrease the step size is: $u_l^2(t) \geq \frac{1}{2}(u_l^*)^2, \forall l \in [L]$. Once this criterion is satisfied, our proof indeed ensures that it would hold for at least up to the early stopping time T_{ub} specified in the theorem. In practice, since u_l^* 's are unknown, we can switch to a more practical criterion: $\max_{l \in [L]} \{|u_l(t+1) - u_l(t)|/|u_l(t) + \varepsilon|\} < \tau$ for some pre-specified tolerance $\tau > 0$ and small value $\varepsilon > 0$ as the criterion for changing the step size. The motivation of this criterion is further discussed in Appendix D. The error bound remains the same as Theorem 3. The change in step size requires a new way to study the gradient dynamics of directions with perturbations. With our proof technique, Theorem 4 requires a smaller bound on δ 's (see Lemma 16 versus Lemma 8 in Appendix C for details). We believe it is a proof artifact and leave the improvement for future work.

Connection to standard sparsity. Consider the degenerate case where each group size is 1. Our reparameterization, together with the normalization step, can roughly be interpreted as $w_i \approx u_i^2 \text{sgn}(v_i)$, which is different from the power-reparameterization $w_i = u_i^N - v_i^N, N \geq 2$ in Vaskevicius et al. (2019) and Li et al. (2021). This also shows why a large step size on v_i is needed at the beginning. If the initialization on v_i is incorrect, the sign of v_i may not move with a small step size.

5 SIMULATION STUDIES

We conduct various experiments on simulated data to support our theory. Following the model in Section 2, we sample the entries of $\mathbf{X}$ i.i.d. using Rademacher random variables and the entries of the noise vector $\boldsymbol{\xi}$ i.i.d. under $N(0, \sigma^2)$. We set $\sigma = 0.5$ throughout the experiments.

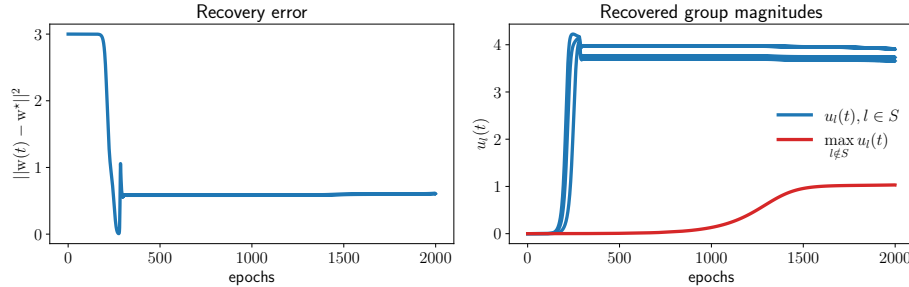


Figure 2: Convergence of Algorithm 1. The entries on the support are all 10.

The effectiveness of our algorithms. We start by demonstrating the convergence of the two proposed algorithms. In this experiment, we set $n = 150$ and $p = 300$. The number of non-zero entries is 9, divided into 3 groups of size 3. We run both Algorithms 1 and 2 with the same initialization $\alpha = 10^{-6}$. The step size γ on $\mathbf{u}$ and decreased step size η on $\mathbf{v}$ are both 10^{-3} . In Figure 2, we present the recovery error of $\mathbf{w}^*$ on the left, and recovered group magnitudes on the right. As we can

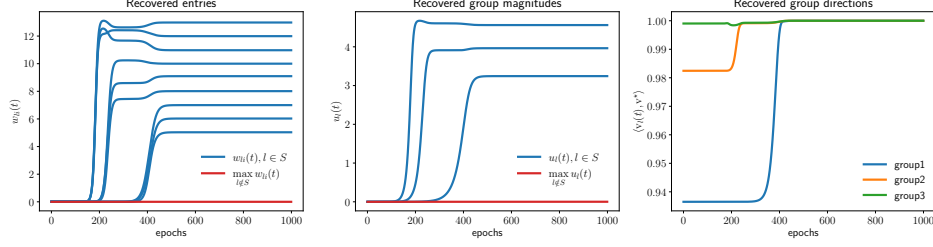
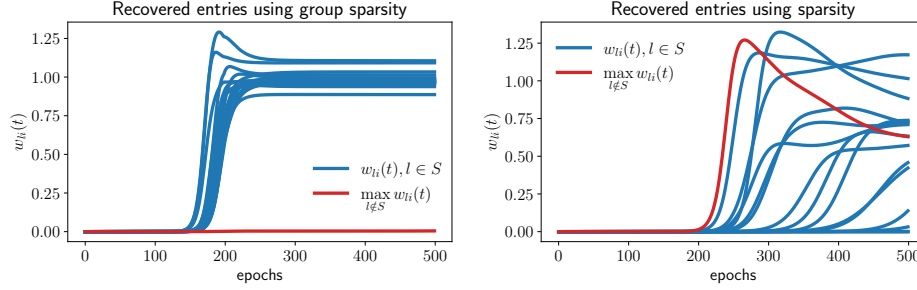
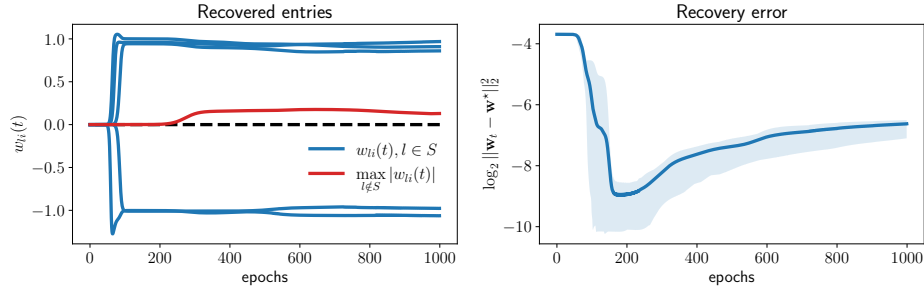


Figure 3: Convergence of Algorithm 2. The entries on the support are from 5 to 13.

see, early stopping is crucial for reaching the structured sparse solution. In Figure 3, we present the recovered entries, recovered group magnitudes and recovered directions for each group from left to right. In addition to convergence, we also observe an incremental learning effect.

Figure 4: Comparison with reparameterization using standard sparsity. $n = 100, p = 500$.

Structured sparsity versus standard sparsity. From our theory, we see that the block incoherence parameter scales with the number of non-zero groups, as opposed to the number of non-zero entries. As such, we can expect an improved sample complexity over the estimators based on unstructured sparse regularization. We choose a larger support size of 16. The entries on the support are all 1 for simplicity. We apply our Algorithm 2 with group size 4. The result is shown in Figure 4 (left). We compare with the method in Vaskevicius et al. (2019) with parameterization $\mathbf{w} = \mathbf{u}^{\odot 2} - \mathbf{v}^{\odot 2}$, designed for unstructured sparsity. We display the result in the right figure, where interestingly, that algorithm fails to converge because of an insufficient number of samples.

Figure 5: Degenerate case when each group size is 1. The $\log \ell_2$ -error plot is repeated 30 times, and the mean is depicted. The shaded area indicates the region between the 25th and 75th percentiles.

Degenerate case. In the degenerate case where each group is of size 1, our reparameterization takes a simpler form $w_i \approx u_i^2 \text{sgn}(v)$, i.e., due to weight normalization, our method normalizes v to 1 or -1 after each step. We demonstrate the efficacy of our algorithms even in the degenerate case. We set $n = 80$ and $p = 200$. The entries on the support are $[1, -1, 1, -1]$ with both positive and negative entries. We present the coordinate plot and the recovery error in Figure 5.

6 DISCUSSION

In this paper, we show that implicit regularization for group-structured sparsity can be obtained by gradient descent (with weight normalization) for a certain, specially designed network architecture. Overall, we hope that such analysis further enhances our understanding of neural network training. Future work includes relaxing the assumptions on δ 's in Theorem 2, and rigorous analysis of modern grouping architectures as well as power parametrizations.

ACKNOWLEDGMENTS

This work was supported in part by the National Science Foundation under grants CCF-1934904, CCF-1815101, and CCF-2005804.

REFERENCES

- Sanjeev Arora, Nadav Cohen, and Elad Hazan. On the optimization of deep networks: Implicit acceleration by overparameterization. In *International Conference on Machine Learning*, pp. 244–253. PMLR, 2018.
- Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo. Implicit regularization in deep matrix factorization. *arXiv preprint arXiv:1905.13655*, 2019.
- Richard G Baraniuk, Volkan Cevher, Marco F Duarte, and Chinmay Hegde. Model-based compressive sensing. *IEEE Transactions on information theory*, 56(4):1982–2001, 2010.
- Raphaël Berthier. Incremental learning in diagonal linear networks. *arXiv preprint arXiv:2208.14673*, 2022.
- Iain Carmichael, Thomas Keefe, Naomi Giertych, and Jonathan P Williams. yaglm: a python package for fitting and tuning generalized linear models that supports structured, adaptive and non-convex penalties. *arXiv preprint arXiv:2110.05567*, 2021.
- Hung-Hsu Chou, Johannes Maly, and Holger Rauhut. More is less: Inducing sparsity via overparameterization. *arXiv preprint arXiv:2112.11027*, 2021.
- Zhen Dai, Mina Karzand, and Nathan Srebro. Representation costs of linear neural networks: Analysis and design. *Advances in Neural Information Processing Systems*, 34, 2021.
- Simon S Du, Wei Hu, and Jason D Lee. Algorithmic regularization in learning deep homogeneous models: Layers are automatically balanced. *Advances in Neural Information Processing Systems*, 31, 2018.
- Yonina C Eldar and Helmut Bolcskei. Block-sparsity: Coherence and efficient recovery. In *2009 IEEE International Conference on Acoustics, Speech and Signal Processing*, pp. 2885–2888. IEEE, 2009.
- Rong Ge, Jason D Lee, and Tengyu Ma. Learning one-hidden-layer neural networks with landscape design. *arXiv preprint arXiv:1711.00501*, 2017.
- Daniel Gissin, Shai Shalev-Shwartz, and Amit Daniely. The implicit bias of depth: How incremental learning drives generalization. *arXiv preprint arXiv:1909.12051*, 2019.
- Suriya Gunasekar, Jason Lee, Daniel Soudry, and Nathan Srebro. Characterizing implicit bias in terms of optimization geometry. In *International Conference on Machine Learning*, pp. 1832–1841. PMLR, 2018a.
- Suriya Gunasekar, Jason Lee, Daniel Soudry, and Nathan Srebro. Implicit bias of gradient descent on linear convolutional networks. *arXiv preprint arXiv:1806.00468*, 2018b.
- Suriya Gunasekar, Jason D Lee, Daniel Soudry, and Nati Srebro. Implicit bias of gradient descent on linear convolutional networks. *Advances in Neural Information Processing Systems*, 31, 2018c.

- Moritz Hardt, Ben Recht, and Yoram Singer. Train faster, generalize better: Stability of stochastic gradient descent. In *International Conference on Machine Learning*, pp. 1225–1234. PMLR, 2016.
- Meena Jagadeesan, Ilya Razenshteyn, and Suriya Gunasekar. Inductive bias of multi-channel linear convolutional networks with bounded weight norm. *arXiv preprint arXiv:2102.12238*, 2021.
- Hui Jin and Guido Montúfar. Implicit bias of gradient descent for mean squared error regression with wide neural networks. *arXiv preprint arXiv:2006.07356*, 2020.
- Li Jing, Jure Zbontar, et al. Implicit rank-minimizing autoencoder. *Advances in Neural Information Processing Systems*, 33:14736–14746, 2020.
- Tae Kwan Lee, Wissam J Baddar, Seong Tae Kim, and Yong Man Ro. Convolution with logarithmic filter groups for efficient shallow cnn. In *International Conference on Multimedia Modeling*, pp. 117–129. Springer, 2018.
- Jiangyuan Li, Thanh Nguyen, Chinmay Hegde, and Raymond K. W. Wong. Implicit sparse regularization: The impact of depth and early stopping. *Advances in Neural Information Processing Systems*, 34, 2021.
- Yuanzhi Li, Tengyu Ma, and Hongyang Zhang. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Conference On Learning Theory*, pp. 2–47. PMLR, 2018.
- Zhiyuan Li, Tianhao Wang, JasonD Lee, and Sanjeev Arora. Implicit bias of gradient descent on reparametrized models: On equivalence to mirror descent. *arXiv preprint arXiv:2207.04036*, 2022.
- Cesare Molinari, Mathurin Massias, Lorenzo Rosasco, and Silvia Villa. Iterative regularization for convex regularizers. In *International conference on artificial intelligence and statistics*, pp. 1684–1692. PMLR, 2021.
- Depen Morwani and Harish G Ramaswamy. Inductive bias of gradient descent for weight normalized smooth homogeneous neural nets. In *International Conference on Algorithmic Learning Theory*, pp. 827–880. PMLR, 2022.
- 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*, pp. 3420–3428. PMLR, 2019.
- Behnam Neyshabur, Ryota Tomioka, Ruslan Salakhutdinov, and Nathan Srebro. Geometry of optimization and implicit regularization in deep learning. *arXiv preprint arXiv:1705.03071*, 2017.
- Mert Pilanci and Tolga Ergen. Neural networks are convex regularizers: Exact polynomial-time convex optimization formulations for two-layer networks. In *International Conference on Machine Learning*, pp. 7695–7705. PMLR, 2020.
- Arda Sahiner, Tolga Ergen, John Pauly, and Mert Pilanci. Vector-output relu neural network problems are copositive programs: Convex analysis of two layer networks and polynomial-time algorithms. *arXiv preprint arXiv:2012.13329*, 2020.
- Tim Salimans and Durk P Kingma. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. *Advances in neural information processing systems*, 29, 2016.
- Todd E Scheetz, Kwang-Youn A Kim, Ruth E Swiderski, Alisdair R Philp, Terry A Braun, Kevin L Knudtson, Anne M Dorrance, Gerald F DiBona, Jian Huang, Thomas L Casavant, et al. Regulation of gene expression in the mammalian eye and its relevance to eye disease. *Proceedings of the National Academy of Sciences*, 103(39):14429–14434, 2006.
- Jonathan Schwarz, Siddhant Jayakumar, Razvan Pascanu, Peter Latham, and Yee Teh. Powerpropagation: A sparsity inducing weight reparameterisation. *Advances in Neural Information Processing Systems*, 34, 2021.
- 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.

- Mihailo Stojnic, Farzad Parvaresh, and Babak Hassibi. On the reconstruction of block-sparse signals with an optimal number of measurements. *IEEE Transactions on Signal Processing*, 57(8): 3075–3085, 2009.
- Twan Van Laarhoven. L2 regularization versus batch and weight normalization. *arXiv preprint arXiv:1706.05350*, 2017.
- Tomas Vaskevicius, Varun Kanade, and Patrick Rebeschini. Implicit regularization for optimal sparse recovery. In *Advances in Neural Information Processing Systems*, pp. 2972–2983, 2019.
- Xiang Wang, Chenwei Wu, Jason D Lee, Tengyu Ma, and Rong Ge. Beyond lazy training for over-parameterized tensor decomposition. *Advances in Neural Information Processing Systems*, 33:21934–21944, 2020.
- Francis Williams, Matthew Trager, Daniele Panozzo, Claudio Silva, Denis Zorin, and Joan Bruna. Gradient dynamics of shallow univariate relu networks. *Advances in neural information processing systems*, 32, 2019.
- 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*, pp. 3635–3673. PMLR, 2020.
- Xiaoxia Wu, Edgar Dobriban, Tongzheng Ren, Shanshan Wu, Zhiyuan Li, Suriya Gunasekar, Rachel Ward, and Qiang Liu. Implicit regularization and convergence for weight normalization. *Advances in Neural Information Processing Systems*, 33:2835–2847, 2020.
- Xuan Wu, Zhijie Zhang, Wanchang Zhang, Yaning Yi, Chuanrong Zhang, and Qiang Xu. A convolutional neural network based on grouping structure for scene classification. *Remote Sensing*, 13(13):2457, 2021.
- Gangcai Xie, Chengliang Dong, Yinfei Kong, Jiang F Zhong, Mingyao Li, and Kai Wang. Group lasso regularized deep learning for cancer prognosis from multi-omics and clinical features. *Genes*, 10(3):240, 2019.
- Saining Xie, Ross Girshick, Piotr Dollár, Zhuowen Tu, and Kaiming He. Aggregated residual transformations for deep neural networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 1492–1500, 2017.
- Yi Yang and Hui Zou. A fast unified algorithm for solving group-lasso penalize learning problems. *Statistics and Computing*, 25:1129–1141, 2015.
- Peng Zhao, Yun Yang, and Qiao-Chu He. Implicit regularization via hadamard product over-parametrization in high-dimensional linear regression. *arXiv preprint arXiv:1903.09367*, 2019.

A GEOMETRIC PROPERTIES OF THE PARAMETRIZATION

We start by calculating the vector field induced by the parameterization $G(\cdot)$.

$$\nabla G_i([\mathbf{u}^\top, \mathbf{v}^\top]) = 2u_{g(i)}v_i\mathbf{e}_{g(i)} + u_{g(i)}^2\mathbf{e}_{L+i},$$

where $\mathbf{e}_i \in \mathbb{R}^{L+p}$ is only 1 on i^{th} entry and 0 elsewhere, and

$$\nabla^2 G_i([\mathbf{u}^\top, \mathbf{v}^\top]) = 2v_i\mathbf{E}_{g(i),g(i)} + 2u_{g(i)}\mathbf{E}_{g(i),L+i} + 2u_{g(i)}\mathbf{E}_{L+i,g(i)},$$

where $\mathbf{E}_{i,j} \in \mathbb{R}^{(L+p) \times (L+p)}$ is the one-hot matrix for i^{th} row and j^{th} column. For $i \neq j$ s.t. $g(i) = g(j)$,

$$\begin{aligned} \nabla^2 G_i([\mathbf{u}^\top, \mathbf{v}^\top])\nabla G_j([\mathbf{u}^\top, \mathbf{v}^\top]) &= (2v_i\mathbf{E}_{g(i),g(i)} + 2u_{g(i)}\mathbf{E}_{g(i),L+i} + 2u_{g(i)}\mathbf{E}_{L+i,g(i)}) \\ &\quad \cdot (2u_{g(j)}v_j\mathbf{e}_{g(j)} + u_{g(j)}^2\mathbf{e}_{L+j}) \\ &= 4u_{g(j)}v_i v_j \mathbf{e}_{g(i)} + 4u_{g(i)}u_{g(j)}v_j \mathbf{e}_{L+i} \\ &= 4u_{g(i)}v_i v_j \mathbf{e}_{g(i)} + 4u_{g(i)}^2 v_j \mathbf{e}_{L+i}, \end{aligned}$$

similarly,

$$\nabla^2 G_j([\mathbf{u}^\top, \mathbf{v}^\top])\nabla G_i([\mathbf{u}^\top, \mathbf{v}^\top]) = 4u_{g(i)}v_i v_j \mathbf{e}_{g(i)} + 4u_{g(i)}^2 v_i \mathbf{e}_{L+j}.$$

Proof for Lemma 1. For two indices within the same group, i.e. $i \neq j$ and $g(i) = g(j)$, we obtain that

$$\begin{aligned} [\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top) &= \nabla^2 G_j([\mathbf{u}^\top, \mathbf{v}^\top])\nabla G_i([\mathbf{u}^\top, \mathbf{v}^\top]) - \nabla^2 G_i([\mathbf{u}^\top, \mathbf{v}^\top])\nabla G_j([\mathbf{u}^\top, \mathbf{v}^\top]) \\ &= 4u_{g(i)}^2 v_j \mathbf{e}_{L+i} - 4u_{g(i)}^2 v_i \mathbf{e}_{L+j}, \end{aligned}$$

which is not always $\mathbf{0}$ when $v_i \neq v_j$. Therefore, $G(\cdot)$ is not commuting. $\square$

Proof for Theorem 1. For $i \neq j$ and $g(i) \neq g(j)$, we have

$$[\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top) = \mathbf{0}.$$

For $i \neq j$ and $g(i) = g(j)$, we have that

$$[\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top) = v_j \nabla G_i - v_i \nabla G_j \in \text{span}\{\nabla G_i\}_{i=1}^p.$$

By Corollary 4.13 in (Li et al., 2022) and Lemma 1, we show that there exists an initialization and a time-dependent loss such that the gradient flow can not be analyzed by mirror flow. $\square$

Alternatively, we can show directly that the necessary condition in Theorem 4.10 in Li et al. (2022) is violated, i.e.,

$$\langle \nabla G_j, [\nabla G_i, [\nabla G_i, \nabla G_j]](\mathbf{u}^\top, \mathbf{v}^\top) \rangle \neq 0$$

for some $[\mathbf{u}^\top, \mathbf{v}^\top]$ in every open set M .

We first obtain that

$$\begin{aligned} \nabla[\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top) &= 8u_{g(i)}v_j\mathbf{E}_{L+i,g(i)} + 4u_{g(i)}^2\mathbf{E}_{L+i,L+j} \\ &\quad - 8u_{g(i)}v_i\mathbf{E}_{L+j,g(i)} - 4u_{g(i)}^2\mathbf{E}_{L+j,L+i}. \end{aligned}$$

Therefore,

$$\begin{aligned} [\nabla G_i, [\nabla G_i, \nabla G_j]](\mathbf{u}^\top, \mathbf{v}^\top) &= \nabla[\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top)\nabla G_i(\mathbf{u}^\top, \mathbf{v}^\top) \\ &\quad - \nabla^2 G_i([\mathbf{u}^\top, \mathbf{v}^\top])[\nabla G_i, \nabla G_j](\mathbf{u}^\top, \mathbf{v}^\top) \\ &= (8u_{g(i)}v_j\mathbf{E}_{L+i,g(i)} + 4u_{g(i)}^2\mathbf{E}_{L+i,L+j} \\ &\quad - 8u_{g(i)}v_i\mathbf{E}_{L+j,g(i)} - 4u_{g(i)}^2\mathbf{E}_{L+j,L+i}) \\ &\quad \cdot (2u_{g(i)}v_i\mathbf{e}_{g(i)} + u_{g(i)}^2\mathbf{e}_{L+i}) \\ &\quad - (2v_i\mathbf{E}_{g(i),g(i)} + 2u_{g(i)}\mathbf{E}_{g(i),L+i} + 2u_{g(i)}\mathbf{E}_{L+i,g(i)}) \\ &\quad \cdot (4u_{g(i)}^2 v_j \mathbf{e}_{L+i} - 4u_{g(i)}^2 v_i \mathbf{e}_{L+j}) \\ &= 16u_{g(i)}^2 v_i v_j \mathbf{e}_{L+i} - 16u_{g(i)}^2 v_i^2 \mathbf{e}_{L+j} - 4u_{g(i)}^4 \mathbf{e}_{L+j} - 8u_{g(i)}^3 v_j \mathbf{e}_{g(i)} \\ &= 16u_{g(i)}^2 v_i v_j \mathbf{e}_{L+i} - (16u_{g(i)}^2 v_i^2 + 4u_{g(i)}^4) \mathbf{e}_{L+j} - 8u_{g(i)}^3 v_j \mathbf{e}_{g(i)}. \end{aligned}$$

Hence,

$$\begin{aligned} & \langle \nabla G_j, [\nabla G_i, [\nabla G_i, \nabla G_j]] \rangle ([\mathbf{u}^\top, \mathbf{v}^\top]) \\ &= \langle 2u_{g(i)}v_j\mathbf{e}_{g(i)} + u_{g(i)}^2\mathbf{e}_{L+j}, 16u_{g(i)}^2v_iv_j\mathbf{e}_{L+i} - (16u_{g(i)}^2v_i^2 + 4u_{g(i)}^4)\mathbf{e}_{L+j} - 8u_{g(i)}^3v_j\mathbf{e}_{g(i)} \rangle \\ &= -16u_{g(i)}^4v_j^2 - 16u_{g(i)}^4v_i^2 - 4u_{g(i)}^6 < 0. \end{aligned}$$

By Theorem 4.10 in Li et al. (2022), there exists an initialization such that no Legendre function R is able to make the gradient flow be written as a mirror flow with respect to R .

B PROOF FOR ANALYSIS OF GRADIENT FLOW

Proof for Lemma 2. Recall

$$\frac{\partial \mathcal{L}}{\partial u_l} = -\frac{2}{n}u_l\mathbf{v}_l^\top \mathbf{X}_l^\top \mathbf{r}(t), \quad \frac{\partial \mathcal{L}}{\partial \mathbf{v}_l} = -\frac{1}{n}u_l^2\mathbf{X}_l^\top \mathbf{r}(t).$$

Therefore, we obtain that

$$\begin{aligned} \frac{\partial \|\mathbf{v}_l(t)\|^2}{\partial t} &= 2\mathbf{v}_l^\top(t) \frac{\partial \mathbf{v}_l(t)}{\partial t} = 2\mathbf{v}_l^\top(t) \left(-\frac{\partial \mathcal{L}}{\partial \mathbf{v}_l} \right) \\ &= \frac{2}{n}u_l^2\mathbf{v}_l^\top(t) \mathbf{X}_l^\top \mathbf{r}(t) \\ &= u_l \left(-\frac{\partial \mathcal{L}}{\partial u_l} \right) = \frac{\partial \frac{1}{2}u_l^2(t)}{\partial t}. \end{aligned}$$

□

Proof for Lemma 3. We start with decomposing $\mathbf{v}_l(0)$

$$\begin{aligned} \mathbf{v}_l(0) &= \eta \frac{1}{n} \mathbf{X}_l^\top \mathbf{y} = \eta \mathbf{w}_l^* + \eta \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X} - \mathbf{I} \right) \mathbf{w}_l^* + \eta \sum_{l' \neq l} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} \mathbf{w}_{l'}^* + \eta \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \\ &= \eta \mathbf{w}_l^* + \eta \mathbf{b}_l. \end{aligned}$$

With this decomposition, we have that

$$\begin{aligned} \langle \mathbf{v}_l(0), \mathbf{v}_l^* \rangle^2 &= \eta^2 ((u_l^*)^2 + \langle \mathbf{b}_l, \mathbf{v}_l^* \rangle)^2 \\ \|\mathbf{v}_l(0)\|_2^2 &= \eta^2 ((u_l^*)^4 + 2\langle \mathbf{b}_l, \mathbf{w}_l^* \rangle + \|\mathbf{b}_l\|_2^2). \end{aligned}$$

Therefore,

$$\begin{aligned} \frac{\langle \mathbf{v}_l(0), \mathbf{v}_l^* \rangle^2}{\|\mathbf{v}_l(0)\|_2^2} &= \frac{\eta^2 ((u_l^*)^2 + \langle \mathbf{b}_l, \mathbf{v}_l^* \rangle)^2}{\eta^2 ((u_l^*)^4 + 2\langle \mathbf{b}_l, \mathbf{w}_l^* \rangle + \|\mathbf{b}_l\|_2^2)} \\ &= 1 - \frac{\|\mathbf{b}_l\|_2^2 - \langle \mathbf{b}_l, \mathbf{v}_l^* \rangle^2}{(u_l^*)^4 + 2\langle \mathbf{b}_l, \mathbf{w}_l^* \rangle + \|\mathbf{b}_l\|_2^2} \\ &= 1 - \frac{\|\mathbf{b}_l/(u_l^*)^2\|_2^2 - \langle \mathbf{b}_l/(u_l^*)^2, \mathbf{v}_l^* \rangle^2}{1 + 2\langle \mathbf{b}_l/(u_l^*)^2, \mathbf{v}_l^* \rangle + \|\mathbf{b}_l/(u_l^*)^2\|_2^2} \\ &= 1 - \frac{1 - \langle \mathbf{b}_l/\|\mathbf{b}_l\|, \mathbf{v}_l^* \rangle^2}{1 + 2\|\mathbf{b}_l\|/(u_l^*)^2 \langle \mathbf{b}_l/\|\mathbf{b}_l\|, \mathbf{v}_l^* \rangle + \|\mathbf{b}_l\|^2/(u_l^*)^4} \|\mathbf{b}_l/(u_l^*)^2\|_2^2 \\ &\geq 1 - \|\mathbf{b}_l/(u_l^*)^2\|_2^2, \end{aligned}$$

where last inequality is from

$$\begin{aligned} \frac{1 - \alpha^2}{\beta^2 + 2\alpha\beta + 1} &= \frac{1}{\frac{\beta^2 + 2\alpha\beta + 1}{1 - \alpha^2}} = \frac{1}{1 + \frac{\beta^2 + 2\alpha\beta + \alpha^2}{1 - \alpha^2}} \\ &= \frac{1}{1 + \frac{(\alpha + \beta)^2}{1 - \alpha^2}} \leq 1, \end{aligned}$$

for $0 \leq \alpha \leq 1$.

Since

$$\|\mathbf{b}_l\|_2 \leq \delta_{in}(u_l^*)^2 + L\delta_{out}(u_l^*)^2 + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_2,$$

we obtain that

$$\left\langle \frac{\mathbf{v}_l(0)}{\|\mathbf{v}_l(0)\|}, \mathbf{v}_l^* \right\rangle \geq 1 - \left(\delta_{in} + L\delta_{out} + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_2 / (u_l^*)^2 \right)^2.$$

□

Lemma 4. Consider a simplified case where $\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l = \mathbf{I}$, $\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} = \mathbf{O}$, $l \neq l'$, if $\mathbf{v}_l(0) = \eta \frac{1}{n} \mathbf{X}_l^\top \mathbf{y}$, then

$$\mathbf{v}_l(t) = c \frac{1}{n} \mathbf{X}_l^\top \mathbf{y},$$

for some constant c .

Proof. From the gradient on the directions, we have that

$$\begin{aligned} \frac{\partial \mathbf{v}_l(t)}{\partial t} &= \frac{1}{n} u_l^2(t) \mathbf{X}_l^\top \mathbf{r}(t) = \frac{1}{n} u_l^2(t) \mathbf{X}_l^\top \mathbf{y} - \frac{1}{n} u_l^2(t) \mathbf{X}_l^\top \sum_{l'} \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \\ &= \frac{1}{n} u_l^2(t) \mathbf{X}_l^\top \mathbf{y} - u_l^4(t) \mathbf{v}_l(t). \end{aligned}$$

Since $\mathbf{v}_l(0)$ is with the same direction as $\frac{1}{n} \mathbf{X}_l^\top \mathbf{y}$ at the initialization. Therefore, $\frac{\partial \mathbf{v}_l(t)}{\partial t}$ has the same direction as $\mathbf{v}_l(t)$. We obtain that $\mathbf{v}_l(t) = c \frac{1}{n} \mathbf{X}_l^\top \mathbf{y}$ for some constant c . □

Lemma 5. If the gradient flow satisfies

$$\frac{1}{2} \frac{\partial u^2(t)}{\partial t} \leq u^6(t) + \sqrt{2} u^4(t) B$$

for some constant $B > 0$, then for any $t \leq T = \frac{\log \frac{1}{\theta}}{2\theta^2 + \theta\sqrt{2}B}$ we have $u(t) \leq \sqrt{\theta}$ with initialization $u(0) = \theta$.

Proof. We wanted to find some time T such that when $t \leq T$, $u(t) \leq \sqrt{\theta}$. Since the gradient is bounded from above, we obtain that

$$\begin{aligned} \frac{1}{2} u^2(T) &\leq \frac{1}{2} \theta^2 \cdot \exp \left(\int_0^T 2u^4(t) + \sqrt{2} u^2(t) B dt \right) \\ &\leq \frac{1}{2} \theta^2 \cdot \exp \left((2\theta^2 + \sqrt{2}\theta B) T \right) \leq \frac{1}{2} \theta. \end{aligned}$$

This gives us

$$T \leq \frac{\log \frac{1}{\theta}}{2\theta^2 + \theta\sqrt{2}B}.$$

□

Lemma 6. Fix any $\tau < \frac{1}{2}$. Consider the gradient flow

$$\frac{1}{2} \frac{\partial u^2(t)}{\partial t} \geq (1 - 2B) \sqrt{2} u^3(t) (u^*)^2 - u^6(t) - \sqrt{2} u^3(t) B (u^*)^2$$

for some constant $0 < B < \frac{1}{10}$ with initialization $u(0) = \theta < \frac{1}{2} u^*$, we have that

$$\left| \frac{1}{\sqrt{2}} u^3(t) - (u^*)^2 \right| < (1 - 3B - \tau) (u^*)^2,$$

after

$$t \geq T = \frac{2^{1/3} (u^*)^{4/3}}{\theta^2} \frac{1}{(1 - 6B) \sqrt{2} (u^*)^2 \theta} + \frac{2 \log_2 \frac{1}{2\tau}}{3 (u^*)^2 (1/2 - 3B) (\sqrt{2} (1/2 - 3B) (u^*)^2)^{1/3}}.$$

Proof. For any $T \geq 0$, we have that

$$\frac{1}{2}u^2(T) \geq \frac{1}{2}\theta^2 \cdot \exp \left(\int_0^T (1-2B)2\sqrt{2}u(t)(u^*)^2 - 2u^4(t) - 2\sqrt{2}u(t)B(u^*)^2 dt \right).$$

When $u(t) < \frac{1}{2}u^*$, we first aim to get T_1 such that $\frac{1}{\sqrt{2}}u^3(T_1) \geq \frac{1}{2}(u^*)^2$. Therefore,

$$\begin{aligned} & \frac{1}{2}\theta^2 \cdot \exp \left(\int_0^T (1-2B)2\sqrt{2}u(t)(u^*)^2 - 2u^4(t) - 2\sqrt{2}u(t)B(u^*)^2 dt \right) \\ & \geq \frac{1}{2}\theta^2 \cdot \exp \left(\left((1-2B)2\sqrt{2}(u^*)^2 - \sqrt{2}(u^*)^2 - 2\sqrt{2}B(u^*)^2 \right) \theta T_1 \right) \\ & \geq \frac{1}{2} \left(\frac{\sqrt{2}}{2}(u^*)^2 \right)^{2/3}. \end{aligned}$$

We obtain that

$$T \geq \frac{2^{1/3}(u^*)^{4/3}}{\theta^2} \frac{1}{(1-6B)\sqrt{2}(u^*)^2\theta}.$$

When $t \geq T_1$, we have that $\frac{1}{\sqrt{2}}u^3(t) \geq \frac{1}{2}(u^*)^2$. Let us denote $\frac{1}{\sqrt{2}}u^3(0) = ((1-3B) - \eta)(u^*)^2$, we wonder how many iterations T_d are needed to make $\frac{1}{\sqrt{2}}u^3(T_d) \geq ((1-3B) - \frac{1}{2}\eta)(u^*)^2$.

$$\begin{aligned} & \frac{1}{2} \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{2/3} \cdot \exp \left(\int_0^T (1-2B)2\sqrt{2}u(t)(u^*)^2 - 2u^4(t) - 2\sqrt{2}u(t)B(u^*)^2 dt \right) \\ & \geq \frac{1}{2} \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{2/3} \cdot \exp \left(\left(\frac{1}{2}\eta(u^*)^2 \right) \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{1/3} T_2 \right) \\ & \geq \frac{1}{2} \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{2/3} \cdot \left(1 + \left(\frac{1}{2}\eta(u^*)^2 \right) \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{1/3} T_2 \right) \\ & \geq \frac{1}{2} \left(\sqrt{2} \left((1-3B) - \frac{1}{2}\eta \right) (u^*)^2 \right)^{2/3}. \end{aligned}$$

Therefore,

$$\begin{aligned} T_2 & \geq \frac{((1-3B) - \frac{1}{2}\eta)^{2/3} - ((1-3B) - \eta)^{2/3}}{((1-3B) - \eta)^{2/3}} \frac{1}{\frac{1}{2}\eta(u^*)^2 \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{1/3}} \\ & \geq \frac{2}{3} \frac{\frac{1}{2}\eta}{\frac{1}{2}\eta(u^*)^2 ((1-3B) - \eta) \left(\sqrt{2}((1-3B) - \eta)(u^*)^2 \right)^{1/3}} \\ & \geq \frac{2}{3(u^*)^2(1/2 - 3B) \left(\sqrt{2}(1/2 - 3B)(u^*)^2 \right)^{1/3}}. \end{aligned}$$

Overall, we obtain that

$$\left| \frac{1}{\sqrt{2}}u^3(t) - (u^*)^2 \right| < (1-3B-\epsilon)(u^*)^2,$$

after

$$t \geq T = T_1 + T_2 \log_2 \frac{1}{2\tau}.$$

□

Proof of Theorem 2. Denote $\zeta = 100 \left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty$. For $l \in S$, the gradient flow can be simplified as

$$\begin{aligned} \frac{1}{2} \frac{\partial u_l^2(t)}{\partial t} &= \frac{2}{n} \mathbf{w}_l^\top(t) \mathbf{X}_l^\top \mathbf{r}(t) \\ &= 2\mathbf{w}_l^\top(t) (\mathbf{w}_l^* - \mathbf{w}_l(t)) + \frac{2}{n} \mathbf{w}_l^\top(t) \mathbf{X}_l^\top \boldsymbol{\xi} \\ &\geq 2u_l^2(t)(u_l^*)^2 \langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle - 2u_l^4(t) \|\mathbf{v}_l(t)\|_2^2 - 2u_l^2(t) \|\mathbf{v}_l(t)\|_2 \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_2. \end{aligned}$$

Since the initialization is balanced $\frac{1}{2}u_l^2(0) = \|\mathbf{v}_l(0)\|_2^2$, we know that from the balancing result Lemma 2,

$$\frac{1}{2}u_l^2(t) = \|\mathbf{v}_l(t)\|_2^2.$$

Since the initialization of $\mathbf{v}_l(t)$ is aligned with direction $\frac{1}{n}\mathbf{X}_l^\top \mathbf{y}$, and with our assumption on orthogonal design, by Lemma 3 and Lemma 4, if $\|\frac{1}{n}\mathbf{X}_l^\top \boldsymbol{\xi}\|_2 \leq B(u_l^*)^2$, we can further simplify the gradient flow as

$$\begin{aligned} \frac{1}{2} \frac{\partial u_l^2(t)}{\partial t} &\geq \sqrt{2}(1 - 2B^2)u_l^3(t)(u_l^*)^2 - u_l^6(t) - \sqrt{2}u_l^3(t)B \\ &\geq \sqrt{2}(1 - 2B)u_l^3(t)(u_l^*)^2 - u_l^6(t) - \sqrt{2}u_l^3(t)B, \end{aligned}$$

where the last inequality holds when $B < 1$. We will verify that $B < 1$ holds in the following analysis.

If $\zeta \geq (u_{max}^*)^2$, then our desired inequality is achieved at the initialization.

If $(u_{min}^*)^2 \leq \zeta \leq (u_{max}^*)^2$, for these group that $\zeta \leq (u_l^*)^2$, applying Lemma 6 with

$$B = \frac{\|\frac{1}{n}\mathbf{X}_l^\top \boldsymbol{\xi}\|_2}{(u_l^*)^2} \leq \frac{\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\|_\infty}{(u_l^*)^2} \leq \frac{1}{100}, \quad \tau = \frac{\epsilon}{(u_l^*)^2}$$

we obtain the convergence on magnitudes

$$|\|\mathbf{w}_l(t)\|_2 - \|\mathbf{w}_l^*\|_2| \leq (3B + \epsilon) \|\mathbf{w}_l^*\|_2,$$

after

$$\frac{2^{1/3}(u_l^*)^{4/3}}{\theta^2} \frac{1}{(1 - 6B)\sqrt{2}(u_l^*)^2\theta} + \frac{2 \log_2 \frac{(u_l)^2}{2\epsilon}}{3(u_l^*)^2(1/2 - 3B) (\sqrt{2}(1/2 - 3B)(u_l^*)^2)^{1/3}}.$$

If $\zeta \leq (u_{min}^*)^2$, similarly applying Lemma 6, the number of iterations needed for entries on the support to converge is

$$T_l = \frac{2^{1/3}(u_{max}^*)^{4/3}}{\theta^2} \frac{1}{(1 - 6B)\sqrt{2}(u_{min}^*)^2\theta} + \frac{2 \log_2 \frac{(u_{max})^2}{2\epsilon}}{3(u_{min}^*)^2(1/2 - 3B) (\sqrt{2}(1/2 - 3B)(u_{min}^*)^2)^{1/3}}.$$

We now have that for $l \in S$,

$$|\|\mathbf{w}_l(t)\|_2 - \|\mathbf{w}_l^*\|_2| \leq (3B + \epsilon) \|\mathbf{w}_l^*\|_2,$$

where $B = \frac{\|\frac{1}{n}\mathbf{X}^\top \mathbf{y}\|_\infty}{(u_{min}^*)^2} \leq \frac{1}{100}, \forall l \in S$.

Recall that the direction is lower bounded by Lemma 3 and Lemma 8,

$$\left\langle \frac{\mathbf{w}_l(t)}{\|\mathbf{w}_l(t)\|_2}, \frac{\mathbf{w}_l^*}{\|\mathbf{w}_l^*\|_2} \right\rangle \geq 1 - B^2.$$

Therefore, the error bound on the support is as follows,

$$\begin{aligned} \|\mathbf{w}_l(t) - \mathbf{w}_l^*\|_\infty &\leq \|\mathbf{w}_l(t) - \mathbf{w}_l^*\|_2 = \left\| \left(\|\mathbf{w}_l(t)\|_2 - (u_l^*)^2 \right) \frac{\mathbf{v}_l(t)}{\|\mathbf{v}_l(t)\|} + (u_l^*)^2 \left\langle \frac{\mathbf{v}_l(t)}{\|\mathbf{v}_l(t)\|}, \mathbf{v}_l^* \right\rangle \right\|_2 \\ &\leq (3B + \tau)(u_l^*)^2 + (u_l^*)^2 \sqrt{2 - 2 \left\langle \frac{\mathbf{v}_l(t)}{\|\mathbf{v}_l(t)\|}, \mathbf{v}_l^* \right\rangle} \\ &= (3B + \tau)(u_l^*)^2 + (u_l^*)^2 \sqrt{2}B \leq \left\| \frac{1}{n}\mathbf{X}^\top \mathbf{y} \right\|_\infty + \epsilon. \end{aligned}$$

For $l \notin S$, we derive a lower bound on the growth rate

$$\begin{aligned} \frac{1}{2} \frac{\partial u_l^2(t)}{\partial t} &= \frac{2}{n} \mathbf{w}_l^\top(t) \mathbf{X}_l^\top \mathbf{r}(t) \\ &= 2 \|\mathbf{w}_l(t)\|_2^2 + \frac{2}{n} \mathbf{w}_l^\top \mathbf{X}_l^\top \boldsymbol{\xi} \\ &\leq u_l^6(t) + \sqrt{2}u_l^4(t)B. \end{aligned}$$

By applying Lemma 5 with $B = \left\| \frac{1}{n} \mathbf{X}^\top \mathbf{y} \right\|_\infty$, we obtain that before

$$T_u = \frac{\log \frac{1}{\theta}}{2\theta^2 + \theta\sqrt{2B}}.$$

Since $\theta < \frac{\epsilon}{2(u_{max})^2}$, $T_l < T_u$ is ensured.

□

C ANALYSIS OF GRADIENT DESCENT

C.1 MONOTONIC UPDATES

Lemma 7. *With an initialization $u(0) < u^*$ and step size $\gamma \leq \frac{1}{4(u^*)^2}$, the updating sequence*

$$u(t) = u(t-1) + 2\gamma u(t-1)[(u^*)^2 - u^2(t-1)],$$

is always bounded above by u^ .*

Proof. We prove it by contradiction. Assume there is a time t s.t.

$$u(t) \leq u^*, u(t+1) > u^*.$$

Therefore,

$$u(t) + 2\gamma u(t)[(u^*)^2 - u^2(t)] > u^*.$$

Denote $\lambda = u(t)/u^*$, we have that

$$1 + 2\gamma(u^*)^2(1 - \lambda^2) - 1/\lambda > 0$$

for some $\lambda \in (0, 1]$.

Let $f(\lambda) = 1 + 2\gamma(u^*)^2(1 - \lambda^2) - 1/\lambda$, we obtain the derivative

$$f'(\lambda) = -4\gamma(u^*)^2\lambda + \frac{1}{\lambda^2} > 0.$$

However, $f_{max}(\lambda) = f(1) = 0$, and $f(\lambda) \leq 0$ for all $\lambda \in (0, 1]$, which gives our desired contradiction. □

C.2 UPDATES WITH BOUNDED PERTURBATIONS

To study the general non-orthogonal and noisy case, we first extend the lemmas above to gradient dynamics with bounded perturbations.

Consider the update on $\mathbf{v}(t)$ with bounded perturbations

$$\begin{aligned} \mathbf{z}(t+1) &= \mathbf{v}(t) + \eta_t u^2(t)((u^*)^2 \mathbf{v}^* - u^2(t) \mathbf{v}(t)) + \eta_t u^2(t) \mathbf{b}_t \\ \mathbf{v}(t+1) &= \frac{\mathbf{z}(t+1)}{\|\mathbf{z}(t+1)\|}. \end{aligned} \quad (4)$$

and the updates on $u(t)$

$$u(t+1) = u(t) + 2\gamma u(t) \mathbf{v}^\top(t+1)\{(u^*)^2 \mathbf{v}^* - u^2(t) \mathbf{v}(t+1)\} + 2\gamma u(t) e_t, \quad (5)$$

Note that if we choose $\eta_t = \frac{1}{u^4(t)}$, Eq. (4) is recast as

$$\begin{aligned} \mathbf{z}(t+1) &= \frac{(u^*)^2}{u^2(t)} \mathbf{v}^* + \frac{1}{u^2(t)} \mathbf{b}_t \\ \mathbf{v}(t+1) &= \frac{\mathbf{z}(t+1)}{\|\mathbf{z}(t+1)\|}. \end{aligned} \quad (6)$$

Lemma 8. Consider the update in Eq. (6), if $\|\mathbf{b}_t\| \leq B(u^*)^2$ for some constant $0 < B < 1$, we have that

$$\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - B^2.$$

Proof. We have that

$$\begin{aligned} \langle \mathbf{z}(t+1), \mathbf{v}^* \rangle &= \frac{(u^*)^2}{u^2(t)} + \frac{1}{u_t^2(t)} \langle \mathbf{b}_t, \mathbf{v}^* \rangle \\ \|\mathbf{z}(t+1)\|^2 &= \frac{(u^*)^4}{u^4(t)} + 2 \frac{(u^*)^2}{u^4(t)} \langle \mathbf{b}_t, \mathbf{v}^* \rangle + \frac{1}{u_t^4(t)} \|\mathbf{b}_t\|^2, \end{aligned}$$

therefore,

$$\begin{aligned} \frac{\langle \mathbf{z}(t+1), \mathbf{v}^* \rangle^2}{\|\mathbf{z}(t+1)\|^2} &= \frac{\frac{(u^*)^4}{u^4(t)} + 2 \frac{(u^*)^2}{u^4(t)} \langle \mathbf{b}_t, \mathbf{v}^* \rangle + \frac{1}{u_t^4(t)} \langle \mathbf{b}_t, \mathbf{v}^* \rangle^2}{\frac{(u^*)^4}{u^4(t)} + 2 \frac{(u^*)^2}{u^4(t)} \langle \mathbf{b}_t, \mathbf{v}^* \rangle + \frac{1}{u_t^4(t)} \|\mathbf{b}_t\|^2} \\ &= 1 - \frac{\|\mathbf{b}_t\|^2 - \langle \mathbf{b}_t, \mathbf{v}^* \rangle^2}{(u^*)^4 + 2(u^*)^2 \langle \mathbf{b}_t, \mathbf{v}^* \rangle + \|\mathbf{b}_t\|^2} \\ &= 1 - \frac{\|\mathbf{b}_t/(u^*)^2\|^2 - \langle \mathbf{b}_t/(u^*)^2, \mathbf{v}^* \rangle^2}{1 + 2 \langle \mathbf{b}_t/(u^*)^2, \mathbf{v}^* \rangle + \|\mathbf{b}_t/(u^*)^2\|^2} \\ &= 1 - \frac{1 - \langle \mathbf{b}_t/\|\mathbf{b}_t\|, \mathbf{v}^* \rangle^2}{1 + 2 \|\mathbf{b}_t\|/(u^*)^2 \langle \mathbf{b}_t/\|\mathbf{b}_t\|, \mathbf{v}^* \rangle + \|\mathbf{b}_t\|^2/(u^*)^4} \|\mathbf{b}_t/(u^*)^2\|^2 \\ &\geq 1 - \|\mathbf{b}_t/(u^*)^2\|^2 \\ &\geq 1 - B^2. \end{aligned}$$

Hence, we have that

$$\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq \sqrt{1 - B^2} \geq 1 - B^2. \quad \square$$

Lemma 9. Consider the updates in Eq. (5) with $|e_t| \leq B$, if $u^2(0) \leq (u^*)^2$, then $u^2(t) \leq (u^*)^2 + B$ for all t . If $u^2(0) \geq (u^*)^2$ and $|\langle \mathbf{v}(t), \mathbf{b}_t \rangle| \leq B_2 \tau (u^*)^2$, then $u^2(t) \geq (1 - B_2)(u^*)^2 - B$ for all t .

Proof. Proof by contradiction similarly to Lemma 7. $\square$

Lemma 10. Fix the step size γ for the update on $u(t)$, and choose $u(0) = \alpha \leq \frac{1}{5}u^*$. Consider the updates in Eq. (5) and Eq. (4) with $|\langle \mathbf{v}(t), \mathbf{b}_t \rangle| \leq \frac{1}{20}(u^*)^2$ and $|e_t| \leq \frac{1}{20}(u^*)^2$, then $T \geq \frac{\log \frac{(u^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2}(u^*)^2)}$, we have that $u^2(T) \geq \frac{1}{2}(u^*)^2$.

Proof. Apply Lemma 8 with $B = \frac{1}{20}$,

$$\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - B^2 = 1 - \frac{1}{400} \geq \frac{4}{5}$$

Starting from $t = 1$, we have that

$$\mathbf{v}^\top(t) \{ (u^*)^2 \mathbf{v}^* - u^2(t) \mathbf{v}(t) \} \geq \frac{4}{5} (u^*)^2 - u^2(t),$$

therefore, we obtain an lower bound of the growth rate on $u(t)$, which reads

$$\begin{aligned} u(t+1) &\geq u(t) + 2\gamma u(t) \left(\frac{4}{5} (u^*)^2 - u^2(t) - \frac{1}{20} (u^*)^2 \right) \\ &= u(t) \left(1 + 2\gamma \left(\frac{3}{4} (u^*)^2 - u^2(t) \right) \right) \\ &\geq u(t) \left(1 + \gamma \frac{1}{2} (u^*)^2 \right). \end{aligned}$$

Therefore, the requirement on the number of iterations is recast as

$$\begin{aligned} \alpha^2 \left(1 + \gamma \frac{1}{2} (u^*)^2\right)^{2T} &\geq \frac{1}{2} (u^*)^2 \\ \iff 2T &\geq \frac{\log \frac{(u^*)^2}{2\alpha^2}}{\log(1 + \gamma \frac{1}{2} (u^*)^2)} \\ \iff T &\geq \frac{\log \frac{(u^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2} (u^*)^2)}. \end{aligned}$$

With these requirements, by Lemma 9, we also have that $u^2(t) \leq \frac{3}{2} (u^*)^2, \forall t \geq 0$. $\square$

Lemma 11. Fix the step size γ for the update on $u(t)$, and choose the initialization $u(0)$ such that $|(u^*)^2 - u^2(0)| \leq \tau (u^*)^2$ where $0 < \tau \leq 1/2$. Consider the updates in Eq. (5) and Eq. (4) with $|\langle \mathbf{v}(t), \mathbf{b}_t \rangle| \leq \frac{1}{10} \tau (u^*)^2$ and $|e_t| \leq \frac{1}{10} \tau (u^*)^2$, then after $T \geq \frac{5}{2\gamma(u^*)^2}$, we have that $\langle \mathbf{v}(t), \mathbf{v}^* \rangle \geq 1 - \frac{1}{5} \tau^2$ for all $t \leq T$ and $|u^2(T) - (u^*)^2| \leq \frac{1}{2} \tau (u^*)^2$.

Proof. When $u^2(0) \leq (u^*)^2$, by applying to Lemma 8, we have that

$$\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - \left(\frac{1}{10} \tau\right)^2 \geq 1 - \frac{1}{5} \tau^2,$$

therefore,

$$\begin{aligned} u(t+1) &\geq u(t) + 2\gamma u(t) \left(\left(1 - \frac{1}{5} \tau\right) (u^*)^2 - u^2(t) - \frac{1}{10} \tau (u^*)^2 \right) \\ &= u(t) \left(1 + 2\gamma \left(\left(1 - \frac{3}{10} \tau\right) (u^*)^2 - u^2(t) \right) \right). \end{aligned}$$

Further, we want to find an lower bound requirement on T s.t.

$$((u^*)^2 - \tau(u^*)^2) \left(1 + 2\gamma \left(\left(1 - \frac{3}{10} \tau\right) (u^*)^2 - \left((u^*)^2 - \frac{1}{2} \tau (u^*)^2 \right) \right) \right)^{2T} \geq (u^*)^2 - \frac{1}{2} \tau (u^*)^2,$$

which can be relaxed as

$$\begin{aligned} ((u^*)^2 - \tau(u^*)^2) \left(1 + \frac{2}{5} \gamma T \tau (u^*)^2 \right) &\geq (u^*)^2 - \frac{1}{2} \tau (u^*)^2 \\ \iff 1 + \frac{2}{5} \gamma T \tau (u^*)^2 &\geq \frac{(u^*)^2 - \frac{1}{2} \tau (u^*)^2}{(u^*)^2 - \tau(u^*)^2} \\ \iff \frac{2}{5} \gamma T \tau (u^*)^2 &\geq \frac{\frac{1}{2} \tau (u^*)^2}{((u^*)^2 - \tau(u^*)^2)} \\ \iff T &\geq \frac{5}{4\gamma(u^*)^2(1-\tau)} \\ \implies T &\geq \frac{5}{2\gamma(u^*)^2}. \end{aligned}$$

When $u^2(0) > (u^*)^2$, we have that

$$\begin{aligned} u(t+1) &\leq u(t) + 2\gamma u(t) \left((u^*)^2 - u^2(t) + \frac{1}{10}\tau(u^*)^2 \right) \\ &= u(t) \left(1 + 2\gamma \left(\left(1 + \frac{1}{10}\tau \right) (u^*)^2 - u^2(t) \right) \right) \\ &\leq u(t) \left(1 - \frac{4}{5}\gamma\tau(u^*)^2 \right). \end{aligned}$$

Similarly, we want to get

$$\begin{aligned} (u^*)^2 + \frac{1}{2}\tau(u^*)^2 &\geq ((u^*)^2 + \tau(u^*)^2) \left(1 - \frac{4}{5}\gamma T\tau(u^*)^2 \right) \\ \iff \frac{(u^*)^2 + \frac{1}{2}\tau(u^*)^2}{(u^*)^2 + \tau(u^*)^2} &\geq 1 - \frac{4}{5}\gamma T\tau(u^*)^2 \\ \iff \frac{4}{5}\gamma T\tau(u^*)^2 &\geq \frac{\frac{1}{2}\tau(u^*)^2}{(u^*)^2 + \tau(u^*)^2} \\ \iff T &\geq \frac{5}{8\gamma(u^*)^2(1+\tau)} \\ \implies T &\geq \frac{5}{8\gamma(u^*)^2}. \end{aligned}$$

If $u(0) \leq u^*$ and $u(t) > u^*$, $t < T$, or $u(0) > u^*$ and $u(t) \leq u^*$, $t < T$, we have already have $|u^2(t) - u^*|^2 \leq \frac{1}{2}\tau(u^*)^2$. By Lemma 9, $|u^2(T) - u^*|^2 \leq \frac{1}{2}\tau(u^*)^2$ remains to hold.

Hence, after $T \geq \frac{5}{2\gamma(u^*)^2}$, we have $|u^2(T) - u^*|^2 \leq \frac{1}{2}\tau(u^*)^2$. $\square$

C.3 ANALYSIS OF PERTURBATIONS

We decompose the updates into several terms for later investigation.

The gradient of $\mathcal{L}(\cdot)$ on each $\mathbf{v}_l$ is

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial \mathbf{v}_l} &= -\frac{1}{n}u_l^2 \mathbf{X}_l^\top \left(\mathbf{y} - \sum_{l' \neq l} u_{l'}^2 \mathbf{X}_{l'} \mathbf{v}_{l'} \right) + \frac{1}{n}u_l^4 \mathbf{X}_l^\top \mathbf{X}_l \mathbf{v}_l \\ &= -\frac{1}{n}u_l^2 \mathbf{X}_l^\top \left(\mathbf{y} - \sum_{l'=1}^L u_{l'}^2 \mathbf{X}_{l'} \mathbf{v}_{l'} \right) \end{aligned}$$

When $l \in S$, the gradient update on each $\mathbf{v}_l$ is

$$\begin{aligned} \mathbf{z}_l(t+1) &= \mathbf{v}_l(t) + \eta_{l,t} u_l^2(t) \frac{1}{n} \mathbf{X}_l^\top \left(\mathbf{y} - \sum_{l'=1}^L u_{l'}^2(t) \mathbf{X}_{l'} \mathbf{v}_{l'}(t) \right) \\ &= \mathbf{v}_l(t) + \eta_{l,t} u_l^2(t) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) \\ &\quad + \eta_{l,t} u_l^2(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) \\ &\quad + \eta_{l,t} u_l^2(t) \sum_{l' \neq l, l' \in S} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \\ &\quad - \eta_{l,t} u_l^2(t) \sum_{l' \in S^c} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \\ &\quad + \eta_{l,t} u_l^2(t) \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi}. \end{aligned}$$

The gradient of $\mathcal{L}(\cdot)$ on each u_l is

$$\begin{aligned}\frac{\partial \mathcal{L}}{\partial u_l} &= -\frac{2}{n} u_l \left\langle \mathbf{X}_l \mathbf{v}_l, \mathbf{y} - \sum_{l' \neq l} u_{l'}^2 \mathbf{X}_{l'} \mathbf{v}_{l'} \right\rangle + \frac{2}{n} u_l^3 \|\mathbf{X}_l \mathbf{v}_l\|^2 \\ &= -\frac{2}{n} u_l \left\langle \mathbf{X}_l \mathbf{v}_l, \mathbf{y} - \sum_{l'=1}^L u_{l'}^2 \mathbf{X}_{l'} \mathbf{v}_{l'} \right\rangle\end{aligned}$$

When $l \in S$, the gradient update on u_l reads

$$\begin{aligned}u_l(t+1) &= u_l(t) + \gamma \frac{2}{n} u_l(t) \left\langle \mathbf{X}_l \mathbf{v}_l(t+1), \mathbf{y} - \sum_{l'=1}^L u_{l'}^2(t) \mathbf{X}_{l'} \mathbf{v}_{l'}(t+1) \right\rangle \\ &= u_l(t) + 2\gamma u_l(t) \mathbf{v}_l^\top(t+1) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t+1)) \\ &\quad + 2\gamma u_l(t) \mathbf{v}_l^\top(t+1) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t+1)) \\ &\quad + 2\gamma u_l(t) \mathbf{v}_l^\top(t+1) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \neq l, l' \in S} \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t+1)) \\ &\quad - 2\gamma u_l(t) \mathbf{v}_l^\top(t+1) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \in S^c} \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t+1) \\ &\quad + 2\gamma u_l(t) \frac{1}{n} \mathbf{v}_l^\top(t+1) \mathbf{X}_l^\top \boldsymbol{\xi}.\end{aligned}$$

We now rewrite the definition of bounded perturbation in Eq. (4, 5), where the bounded perturbation $e_{l,t}$ on updates of $u_l(t)$ reads

$$\begin{aligned}e_{l,t} &= \mathbf{v}_l^\top(t+1) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t+1)) \\ &\quad + \mathbf{v}_l^\top(t+1) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \neq l, l' \in S} \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t+1)) \\ &\quad - \mathbf{v}_l^\top(t+1) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \in S^c} \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t+1) \\ &\quad + \frac{1}{n} \mathbf{v}_l^\top(t+1) \mathbf{X}_l^\top \boldsymbol{\xi},\end{aligned}$$

and the bounded perturbation $\mathbf{b}_{l,t}$ on updates of $\mathbf{v}_l(t)$ reads

$$\begin{aligned}\mathbf{b}_{l,t} &= \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) \\ &\quad + \sum_{l' \neq l, l' \in S} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \\ &\quad - \sum_{l' \in S^c} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \\ &\quad + \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi}.\end{aligned}$$

We show in Lemma 11 that when the perturbations are bounded, the direction is roughly accurate ($\langle \mathbf{v}_l(t), \mathbf{v}^* \rangle$ is large) and $u_l(t)$ converges exponentially. Now we show below that when the direction is roughly accurate and $u_l(t)$ is close to u_l^* , the perturbations are bounded.

Lemma 12. Assume $\delta_{in} \leq \frac{(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{(u_{min}^*)^2}{120s(u_{max}^*)^2}$, $\alpha < \frac{1}{2}\sqrt{\frac{\tau_0}{L}}u_l^*$, $\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\|_\infty \leq \frac{1}{80}\tau_0(u_l^*)^2$ and $|(u_l^*)^2 - u_l^2(0)| \leq \tau(u_l^*)^2$ for each $l \in [L]$ where $0 < \tau_0 \leq \tau \leq 1/2$. If $\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle \geq 1 - \frac{1}{5}\tau^2$, then $|\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{10}\tau(u_l^*)^2$ and $|e_{l,t}| \leq \frac{1}{10}\tau(u_l^*)^2$.

Proof. We first verify

$$\begin{aligned} \|(u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)\| &= \|\{(u_l^*)^2 - u_l^2(t)\} \mathbf{v}_l^* - u_l^2(t) \{\mathbf{v}_l(t) - \mathbf{v}_l^*\}\| \\ &\leq |(u_l^*)^2 - u_l^2(t)| + u_l^2(t) \|\mathbf{v}_l(t) - \mathbf{v}_l^*\| \\ &\leq \tau(u_l^*)^2 + u_l^2(t) \sqrt{2 - 2\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle} \\ &\leq \tau(u_l^*)^2 + \frac{3}{2}(u_l^*)^2 \frac{\sqrt{2}}{\sqrt{5}}\tau \\ &\leq 3\tau(u_l^*)^2. \end{aligned} \tag{7}$$

By Assumption 1, we have that

$$\begin{aligned} &\left| \mathbf{v}_l^\top(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) + \mathbf{v}_l^\top(t) \sum_{l' \neq l, l' \in S} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \right| \\ &\leq 3\delta_{in}\tau(u_{max}^*)^2 + 3s\delta_{out}\tau(u_{max}^*)^2 \leq \frac{1}{40}\tau(u_l^*)^2 + \frac{1}{40}\tau(u_l^*)^2 = \frac{1}{20}\tau(u_l^*)^2. \end{aligned}$$

For the other two terms, we have that

$$\left| \mathbf{v}_l^\top(t) \sum_{l' \in S^c} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \right| \leq \delta(L-s)\alpha^2 \leq \frac{1}{80}\tau(u_l^*)^2,$$

and

$$\begin{aligned} \left| \mathbf{v}_l^\top(t) \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right| &\leq \|\mathbf{v}_l(t)\|_1 \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty \\ &\leq \|\mathbf{v}_l(t)\|_2 \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty \\ &\leq \frac{1}{80}\tau(u_l^*)^2. \end{aligned}$$

Therefore,

$$|e_{l,t}| = |\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{20}\tau(u_l^*)^2 + \frac{1}{80}\tau(u_l^*)^2 + \frac{1}{80}\tau(u_l^*)^2 \leq \frac{1}{10}\tau(u_l^*)^2.$$

□

Lemma 11 shows that when the upper bound of perturbation is fixed, $u_l(t)$ grows. Now we show that after $u_l(t)$ grows, the upper bound of perturbations will be decreased.

Lemma 13. Assume $\delta_{in} \leq \frac{(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{(u_{min}^*)^2}{120s(u_{max}^*)^2}$, $\alpha < \frac{\sqrt{\tau_0}}{2\sqrt{L}}u_l^*$, $\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\|_\infty \leq \frac{1}{80}\tau_0(u_l^*)^2$ and $\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle \geq 1 - \frac{1}{5}\tau^2$. If we achieve that $|(u_l^*)^2 - u_l^2(0)| \leq \frac{1}{2}\tau(u_l^*)^2$ for each $l \in [L]$ where $0 < \tau_0 \leq \tau \leq 1/2$, then $|\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{20}\tau(u_l^*)^2$ and $|e_{l,t}| \leq \frac{1}{20}\tau(u_l^*)^2$.

Proof. Similarly to the proof of Lemma 11,

$$\begin{aligned} \|(u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)\| &\leq \frac{1}{2}\tau(u_l^*)^2 + u_l^2(t) \sqrt{2 - 2\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle} \\ &\leq \frac{1}{2}\tau(u_l^*)^2 + \frac{3}{2}(u_l^*)^2 \frac{1}{\sqrt{5}}\tau \\ &\leq \frac{3}{2}\tau(u_l^*)^2. \end{aligned}$$

By Assumption 1, we have that

$$\begin{aligned} & \left| \mathbf{v}_l^\top(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) + \mathbf{v}_l^\top(t) \sum_{l' \neq l, l' \in S} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \right| \\ & \leq \frac{3}{2} \delta_{in} \tau (u_{max}^*)^2 + \frac{3}{2} s \delta_{out} \tau (u_{max}^*)^2 \leq \frac{1}{40} \tau (u_l^*)^2, \end{aligned}$$

where $\delta \leq \frac{1}{60s}$. Similarly, we obtain that

$$|e_{l,t}| = |\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{40} \tau (u_l^*)^2 + \frac{1}{80} \tau (u_l^*)^2 + \frac{1}{80} \tau (u_l^*)^2 \leq \frac{1}{20} \tau (u_l^*)^2.$$

□

By Lemma 10, we know that after certain iterations, we have that $|u^2(t) - (u^*)^2| \leq \frac{1}{2} (u^*)^2$. Starting from there, we will apply Lemma 11 and Lemma 12 iteratively until we have our desired accuracy.

We just need to verify when $\tau = \frac{1}{2}$, the condition of either Lemma 11 and Lemma 12 is satisfied. Note that the condition of Lemma 10 already satisfies the condition of Lemma 11 at $\tau = \frac{1}{2}$. Note the condition of Lemma 10 is satisfied when $\delta_{in} \leq \frac{(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{(u_{min}^*)^2}{120s(u_{max}^*)^2}$, $\alpha \leq \frac{1}{4} (u_{min}^*)^2$, $\left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \leq \frac{1}{80} \tau_0 (u_{min}^*)^2$.

C.4 ERROR ANALYSIS OUTSIDE THE SUPPORT

We only care about the growth rate of $u_l(t)$ when $l \notin S$. When $l \in S^c$, the gradient updates on u_l reads

$$\begin{aligned} u_l(t+1) &= u_l(t) + \gamma \frac{2}{n} u_l(t) \left\langle \mathbf{X}_l \mathbf{v}_l(t), \mathbf{y} - \sum_{l'=1}^L u_{l'}^2(t) \mathbf{X}_{l'} \mathbf{v}_{l'}(t) \right\rangle \\ &= u_l(t) - 2\gamma u_l^3(t) \\ &\quad - 2\gamma u_l^3(t) \mathbf{v}_l^\top(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) \mathbf{v}_l(t) \\ &\quad + 2\gamma u_l(t) \mathbf{v}_l^\top(t) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \in S} \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \\ &\quad - 2\gamma u_l(t) \mathbf{v}_l^\top(t) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \neq l, l' \in S^c} \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \\ &\quad + 2\gamma u_l(t) \frac{1}{n} \mathbf{v}_l(t) \mathbf{X}_l^\top \boldsymbol{\xi}. \end{aligned}$$

Consider the initialization is $u_l(0) = \alpha$, we wonder the smallest number t of iterations that we can ensure $u_l(t) \leq \sqrt{\alpha}$. Denote

$$\begin{aligned} e_{l,t} &= -u_l^2(t) - u_l^2(t) \mathbf{v}_l^\top(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) \mathbf{v}_l(t) \\ &\quad + \mathbf{v}_l^\top(t) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \in S} \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \\ &\quad - \mathbf{v}_l^\top(t) \frac{1}{n} \mathbf{X}_l^\top \sum_{l' \neq l, l' \in S^c} \mathbf{X}_{l'} u_{l'}^2(t) \mathbf{v}_{l'}(t) \\ &\quad + \frac{1}{n} \mathbf{v}_l^\top(t) \mathbf{X}_l^\top \boldsymbol{\xi}. \end{aligned}$$

We have that

$$|e_{l,t}| \leq \alpha + \alpha \delta_{in} + \alpha \delta_{out} (L - s) + \frac{3}{2} (u_{max}^*)^2 \delta_{out} s + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty.$$

If $\alpha \leq \frac{1}{80L}(u_{min}^*)^2$, $\delta_{in} \leq \frac{(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{(u_{min}^*)^2}{120s(u_{max}^*)^2}$, we have that

$$|e_{l,t}| \leq \frac{1}{20}(u_{min}^*)^2 + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty. \quad (8)$$

Lemma 14. *Consider*

$$u(t+1) = u(t)(1 + 2\gamma e_t)$$

where $|e_t| \leq B$ and $u(0) = \alpha$. Let the step size $\gamma \leq \frac{1}{4B}$, then for any $t \leq T = \frac{1}{32\gamma B} \log \frac{1}{\alpha^4}$, we have $u(t) \leq \sqrt{u(0)}$.

Proof. We start by observing,

$$\begin{aligned} \sqrt{\alpha} &\geq u(t) \geq \alpha(1 + 2\gamma B)^t \\ \Leftrightarrow t &\leq \frac{\log \frac{1}{\sqrt{\alpha}}}{\log(1 + 2\gamma B)}. \end{aligned}$$

By using $\log x \leq x - 1$,

$$\frac{\log \frac{1}{\sqrt{\alpha}}}{\log(1 + 2\gamma B)} \geq \frac{1}{2\gamma B} \log \frac{1}{\sqrt{\alpha}} \geq \frac{1}{32\gamma B} \log \frac{1}{\alpha^4}.$$

□

D PROOF FOR THEOREMS IN SECTION 4

D.1 PROOF OF THEOREM 3

Proof. If $\zeta \geq (u_{max}^*)^2$, at the initialization, we already have for $\forall l \in [L]$

$$\begin{aligned} \|u_l^2(0)\mathbf{v}_l(0) - (u_l^*)^2\mathbf{v}_l^*\|_\infty &\leq u_l^2(0) + (u_l^*)^2 \leq \alpha^2 + (u_{max}^*)^2 \\ &\leq 2(u_{max}^*)^2 \leq 2\zeta \\ &\leq 160 \left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \vee 160\epsilon. \end{aligned}$$

If $\zeta \leq (u_{max}^*)^2$, for those $l \in S$ such that $\zeta \leq (u_l^*)^2$, we can apply Lemma 10. After

$$T_1 = \frac{\log \frac{(u_l^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2}(u_l^*)^2)},$$

we obtain that $\frac{1}{2}(u_l^*)^2 \leq u_l^2(T_1) \leq \frac{3}{2}(u_l^*)^2$, where we also have that $\left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \leq \frac{1}{80}(u_l^*)^2$ for every l .

Let m_0 be the number s.t.

$$2^{-m_0-1}(u_{max}^*)^2 \leq \zeta \leq 2^{-m_0}(u_{max}^*)^2,$$

which can be written as $m_0 = \lfloor \log_2 \frac{(u_{max}^*)^2}{\zeta} \rfloor$. We can apply Lemma 11 and Lemma 12 together m_0 times. Then further after

$$T_2 = \lfloor \log_2 \frac{(u_{max}^*)^2}{\zeta} \rfloor \frac{5}{2\gamma(u_l^*)^2},$$

we have that

$$\begin{aligned} |u_l^2(T_2) - (u_l^*)^2| &\leq 2^{-m_0}(u_{max}^*)^2 \leq 2\zeta \\ \langle \mathbf{v}_l(T_2), \mathbf{v}_l^* \rangle &\geq 1 - \frac{1}{5}2^{-2m_0}. \end{aligned}$$

Therefore,

$$\begin{aligned}
\|u_l^2(T_2)\mathbf{v}_l(T_2) - (u_l^*)^2\mathbf{v}_l^*\|_\infty &\leq \|u_l^2(T_2)\mathbf{v}_l(T_2) - (u_l^*)^2\mathbf{v}_l^*\|_2 \\
&\leq \|(u_l^2(T_2) - (u_l^*)^2)\mathbf{v}_l(T_2) - (u_l^*)^2(\mathbf{v}_l^* - \mathbf{v}_l(T_2))\|_2 \\
&\leq 2^{-m_0}(u_{max}^*)^2 + (u_l^*)^2\sqrt{2 - 2\langle \mathbf{v}_l(T_2), \mathbf{v}_l^* \rangle} \\
&\leq 2^{-m_0}(u_{max}^*)^2 + (u_l^*)^2\frac{2}{5}2^{-m_0} \\
&\leq 2\zeta.
\end{aligned} \tag{9}$$

Note that the above inequality holds for every $l \in S$ such that $(u_l^*)^2 \geq \zeta$. For those l such that $\zeta \geq (u_l^*)^2$, we are not able to recover the true signal $(u_l^*)^2$. the gradient dynamics on this group behaves as errors outside group, and bounded by Lemma 14.

For entries outside the support, we know that from Eq. (8),

$$B = \frac{1}{20}(u_{min}^*)^2 + \left\| \frac{1}{n} \mathbf{X}_l^\top \boldsymbol{\xi} \right\|_\infty \leq \frac{1}{10}(\zeta \vee (u_{min}^*)^2).$$

By Lemma 14, we have that before $T_3 \leq \frac{1}{32\gamma B} \log \frac{1}{\alpha^4}$, $u_l(T_3) \leq \sqrt{\alpha}$.

When $\zeta \leq (u_{min}^*)^2$, Eq. (9) holds for every $l \in S$. Therefore, a uniform number of iterations T_1 and T_2 for all groups is written as

$$T_1 = \frac{\log \frac{(u_{max}^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2}(\zeta \vee (u_{min}^*)^2))},$$

and

$$T_2 = \lfloor \log_2 \frac{(u_{max}^*)^2}{\zeta} \rfloor \frac{5}{2\gamma(\zeta \vee (u_{min}^*)^2)}.$$

All we left is to show that $T_3 \geq T_1 + T_2$. We observe that

$$\begin{aligned}
T_1 &= \frac{\log \frac{(u_{max}^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2}(\zeta \vee (u_{min}^*)^2))} \leq \frac{1 + \gamma \frac{1}{2}(\zeta \vee (u_{min}^*)^2)}{\gamma(\zeta \vee (u_{min}^*)^2)} \log \frac{(u_{max}^*)^2}{2\alpha^2} \\
&\leq \frac{2}{\gamma(\zeta \vee (u_{min}^*)^2)} \log \frac{(u_{max}^*)^2}{2\alpha^2}
\end{aligned}$$

where the first inequality is by $\log x \geq \frac{x-1}{x}$.

With our choice of small initialization on α , we have $T_1 \leq \frac{1}{2}T_3$, due to $\alpha < \frac{1}{(u_{max}^*)^8}$. We have $T_2 \leq \frac{1}{2}T_3$, because of $\alpha < \frac{\zeta^4}{(u_{max}^*)^8}$.

Hence, we obtain that after $T_l = T_1 + T_2 \geq \frac{\log \frac{(u_{max}^*)^2}{2\alpha^2}}{2 \log(1 + \gamma \frac{1}{2}(\zeta \vee (u_{min}^*)^2))} + \lfloor \log_2 \frac{(u_{max}^*)^2}{\zeta} \rfloor \frac{5}{2\gamma(\zeta \vee (u_{min}^*)^2)}$, and before $T_u = T_3 \leq \frac{5}{16\gamma(\zeta \vee (u_{min}^*)^2)} \log \frac{1}{\alpha^4}$,

$$\|u_l^2(t)\mathbf{v}_l(t) - (u_l^*)^2\mathbf{v}_l^*\|_\infty \lesssim \begin{cases} \left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \vee \epsilon, & \text{if } l \in S. \\ \alpha, & \text{if } l \notin S. \end{cases}$$

□

D.2 PROOF FOR COROLLARY 1

Here is a standard result for sub-Gaussian noise.

Lemma 15. Let $\frac{1}{\sqrt{n}}\mathbf{X}$ be a $n \times p$ matrix with ℓ_2 -normalized columns. Let $\boldsymbol{\xi} \in \mathbb{R}^n$ be a vector of independent σ^2 -sub-Gaussian random variables. Then, with probability at least $1 - \frac{1}{8p^3}$

$$\left\| \frac{1}{n} \mathbf{X}^\top \boldsymbol{\xi} \right\|_\infty \lesssim \sqrt{\frac{\sigma^2 \log p}{n}}.$$

Proof of Lemma 15. Since the vector $\boldsymbol{\xi}$ are made of independent σ^2 -sub-Gaussian random variables and any column of $\mathbf{X}$ is ℓ_2 -normalized, the random variable $\frac{1}{\sqrt{n}}(\mathbf{X}^\top \boldsymbol{\xi})_i$ is still σ^2 -sub-Gaussian.

It is a standard result that for any $\epsilon > 0$,

$$\mathbb{P}\left(\left\|\frac{1}{\sqrt{n}}\mathbf{X}^\top \boldsymbol{\xi}\right\|_\infty > \epsilon\right) \leq 2p \exp\left(-\frac{\epsilon^2}{2\sigma^2}\right).$$

Setting $\epsilon = 2\sqrt{2\sigma^2 \log(2p)}$, with probability at least $1 - \frac{1}{8p^3}$ we have

$$\left\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\right\|_\infty \leq \frac{1}{\sqrt{n}} 2\sqrt{\sigma^2 \log(2p)} \lesssim \sqrt{\frac{\sigma^2 \log p}{n}}.$$

□

Proof of Corollary 1. Since $\boldsymbol{\xi}$ is made of independent σ^2 -sub-Gaussian entries, by Lemma 15 with probability $1 - 1/(8p^3)$ we have

$$\left\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\right\|_\infty \leq 2\sqrt{\frac{2\sigma^2 \log(2p)}{n}}.$$

Hence, letting $\epsilon = 2\sqrt{\frac{2\sigma^2 \log(2p)}{n}}$, we obtain that

$$\|(\mathbf{D}\mathbf{u}(t))^2 \odot \mathbf{v}(t) - \mathbf{w}^*\|_2^2 \lesssim \sum_{l \in S} \epsilon^2 + \sum_{l \notin S} \alpha \leq s\epsilon^2 + (L-s)\frac{\epsilon^2}{L^2} \lesssim \frac{s\sigma^2 \log p}{n}.$$

□

D.3 CONVERGENCE FOR ALGORITHM 2

Lemma 16. Consider the update in Eq. (4), choose the step size $\eta_t = \eta \leq \frac{4}{9(u^*)^4}$, if $\langle \mathbf{v}(t), \mathbf{v}^* \rangle \geq 1 - \frac{1}{5}\tau$, $|u^2(t) - (u^*)^2| \leq \tau(u^*)^2$ and $\|\mathbf{b}_t\| \leq \frac{1}{10}\tau(u^*)^2$ for some constant $0 < \tau < \frac{1}{2}$, we have that

$$\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - \frac{1}{5}\tau.$$

Proof. We first rewrite $\mathbf{z}(t+1)$ as

$$\mathbf{z}(t+1) = \eta u^2(t)(u^*)^2 \mathbf{v}^* + (1 - \eta u^4(t))\mathbf{v}(t) + \eta u^2(t)\mathbf{b}_t.$$

Therefore,

$$\begin{aligned} \langle \mathbf{z}(t+1), \mathbf{v}^* \rangle &\geq \eta u^2(t)(u^*)^2 + (1 - \eta u^4(t))\langle \mathbf{v}(t), \mathbf{v}^* \rangle + \eta u^2(t)\langle \mathbf{b}_t, \mathbf{v}^* \rangle \\ &\geq \eta u^2(t)(u^*)^2 + (1 - \eta u^4(t))\left(1 - \frac{1}{5}\tau\right) - \eta u^2(t)\frac{1}{10}\tau(u^*)^2 \\ \|\mathbf{z}(t+1)\| &\leq \eta u^2(t)(u^*)^2 + (1 - \eta u^4(t)) + \eta u^2(t)\frac{1}{10}\tau(u^*)^2. \end{aligned}$$

We obtain that

$$\begin{aligned} \langle \mathbf{v}(t+1), \mathbf{v}^* \rangle &= \frac{\langle \mathbf{z}(t+1), \mathbf{v}^* \rangle}{\|\mathbf{z}(t+1)\|} \geq 1 - \frac{\frac{1}{5}\tau(1 - \eta u^4(t)) + 2\eta u^2(t)\frac{1}{10}\tau(u^*)^2}{\eta u^2(t)(u^*)^2 + (1 - \eta u^4(t)) + \eta u^2(t)\frac{1}{10}\tau(u^*)^2} \\ &\geq 1 - \frac{1 - \eta u^4(t) + \eta u^2(t)(u^*)^2}{\eta u^2(t)(u^*)^2 + (1 - \eta u^4(t)) + \eta u^2(t)\frac{1}{10}\tau(u^*)^2} \frac{1}{5}\tau \\ &\geq 1 - \frac{1}{5}\tau. \end{aligned}$$

□

Note that compared with Lemma 8, under the condition $\|\mathbf{b}_t\| \leq B(u^*)^2$, we get $\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - B$ instead of $\langle \mathbf{v}(t+1), \mathbf{v}^* \rangle \geq 1 - B^2$. Accordingly, we need a new version for Lemma 12 with a smaller bound on δ to make up the loss in Lemma 16.

Lemma 17. Assume $\delta_{in} \leq \frac{\sqrt{\tau_0}(u_{min}^*)^2}{120(u_{max}^*)^2}$ and $\delta_{out} \leq \frac{\sqrt{\tau_0}(u_{min}^*)^2}{120s(u_{max}^*)^2}$, $\alpha < \frac{1}{2}\sqrt{\frac{\tau_0}{L}}u_l^*$, $\|\frac{1}{n}\mathbf{X}^\top \boldsymbol{\xi}\|_\infty \leq \frac{1}{80}\tau_0(u_l^*)^2$ and $|(u_l^*)^2 - u_l^2(0)| \leq \tau(u_l^*)^2$ for each $l \in [L]$ where $0 < \tau_0 \leq \tau \leq 1/2$. If $\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle \geq 1 - \frac{1}{5}\tau$, then $|\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{10}\tau(u_l^*)^2$ and $|e_{l,t}| \leq \frac{1}{10}\tau(u_l^*)^2$.

Proof. Similarly to Lemma 12, we have that

$$\begin{aligned} \|(u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)\| &\leq \tau(u_l^*)^2 + u_l^2(t) \sqrt{2 - 2\langle \mathbf{v}_l(t), \mathbf{v}_l^* \rangle} \\ &\leq \tau(u_l^*)^2 + \frac{3}{2}(u_l^*)^2 \frac{\sqrt{2}}{\sqrt{5}} \sqrt{\tau} \\ &\leq \left(1 + 2\frac{1}{\sqrt{\tau_0}}\right) \tau(u_l^*)^2. \end{aligned} \quad (10)$$

By Assumption 1, we have that

$$\begin{aligned} &\left| \mathbf{v}_l^\top(t) \left(\frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_l - \mathbf{I} \right) ((u_l^*)^2 \mathbf{v}_l^* - u_l^2(t) \mathbf{v}_l(t)) + \mathbf{v}_l^\top(t) \sum_{l' \neq l, l' \in S} \frac{1}{n} \mathbf{X}_l^\top \mathbf{X}_{l'} ((u_{l'}^*)^2 \mathbf{v}_{l'}^* - u_{l'}^2(t) \mathbf{v}_{l'}(t)) \right| \\ &\leq \left(1 + 2\frac{1}{\sqrt{\tau_0}}\right) \delta_{in} \tau(u_{max}^*)^2 + \left(1 + 2\frac{1}{\sqrt{\tau_0}}\right) s \delta_{out} \tau(u_{max}^*)^2 \leq \frac{1}{20} \tau(u_l^*)^2, \end{aligned}$$

where $\delta \leq \frac{\sqrt{\tau_0}(u_{min}^*)^2}{60s(u_{max}^*)^2}$. The other two terms follows exactly what we did in Lemma 12. Therefore,

$$|e_{l,t}| = |\langle \mathbf{v}_l(t), \mathbf{b}_{l,t} \rangle| \leq \frac{1}{20} \tau(u_l^*)^2 + \frac{1}{80} \tau(u_l^*)^2 + \frac{1}{80} \tau(u_l^*)^2 \leq \frac{1}{10} \tau(u_l^*)^2.$$

□

Proof to Theorem 4. The proof is similar to that of Theorem 3. For the first stage, we apply Lemma 10, as nothing is changed from Theorem 3. For the second stage, instead of applying Lemma 11 and Lemma 12, we apply Lemma 16 and Lemma 17 iteratively. To apply these lemmas, we first observe that

$$\zeta \leq \tau_0(u_{max}^*)^2 \iff \frac{\zeta}{(u_{max}^*)^2} \leq \tau_0.$$

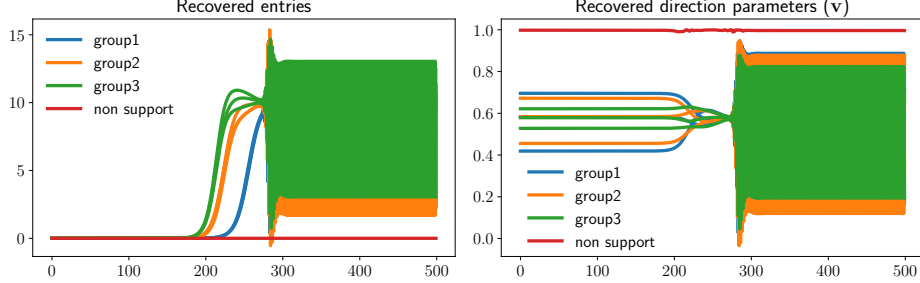
Therefore the requirement on δ 's becomes $\delta_{in} \leq \frac{\sqrt{\tau_0}(u_{min}^*)^2}{120(u_{max}^*)^3}$ and $\delta_{out} \leq \frac{\sqrt{\tau_0}(u_{min}^*)^2}{120s(u_{max}^*)^3}$. The number of iterations and convergence results follow from the proof of Theorem 3.

□

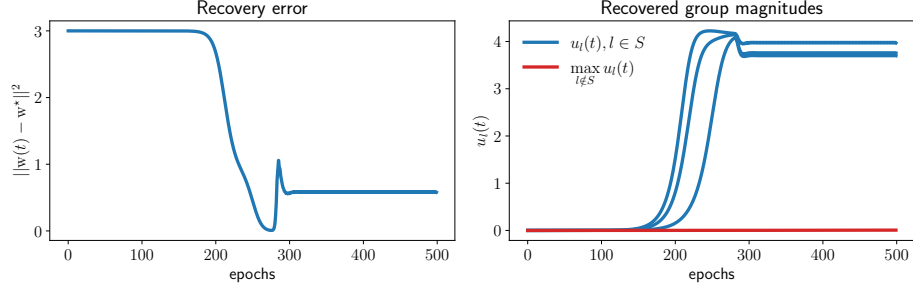
The criterion for switching time. We provide some motivation for the practical criterion. We first note that, the criterion in Theorem 4 actually indicates a lower bound of switching time. With more derivations, our results still hold if one choose to switch after the time when the criterion is first satisfied (instead of switching right at that time.) Let us focus on the entries on the support. In the proof of Theorem 3, one can also obtain the convergence on $u_l(t)$ as the positiveness of $u_l(t)$ can be ensured with a small step size γ (since the power-parametrization will recast the gradient updates into a multiplicative sequence). Therefore, with an appropriate choice of τ , the practical criterion $\max_{l \in S} \{ |u_l(t+1) - u_l(t)| / |u_l(t) + \varepsilon| \} < \tau$ would imply the theoretical criterion $u_l(t)^2 \geq \frac{1}{2} u_l^*(t)^2$ on the support, and therefore would indicate a possibly later switching time than what the theoretical criterion determines. For gradient updates outside the support, we observe slow growth rate and hence the practical rule is likely satisfied on the non-support entry, which we observe in the numerical experiments. Note that the switching only happens when both the support and non-support entries fulfill the criterion.

E MORE NUMERICAL RESULTS

E.1 STABILITY ISSUE OF ALGORITHM 1 AND STANDARD GD



(a) Numerical instability in direction estimations.



(b) Parameter estimation error remains small.

Figure 6: Numerical instability of algorithm 1

Stability issue of Algorithm 1. Figure 6 presents the recovered entries and direction parameters $\mathbf{v}(t)$ under the same setting as Figure 2. Because of the large learning rate on $\mathbf{v}$, the algorithm may not show a convergent result in the latter stage due to the irreducible error (perturbations). Although the parameter estimation is still reasonable with normalization on each $\mathbf{v}_l, l \in [L]$, we still aim to get a stable algorithm, which motivates our algorithm 2.

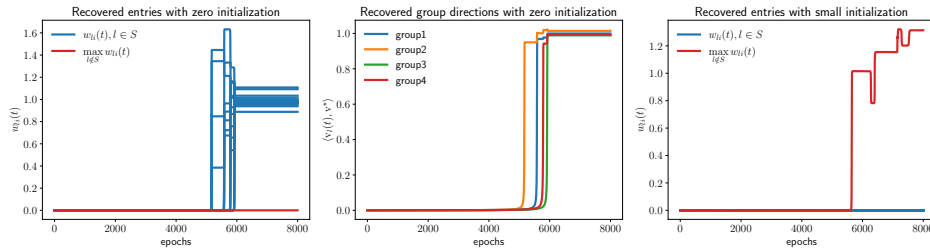


Figure 7: Gradient descent without weight normalization.

Standard gradient descent. To further understand how weight normalization affects the gradient dynamics, we conduct experiments using standard gradient descent without weight normalization. For that, we use the same setting as in Figure 4 and show the result in Figure 7. The left and middle figures are based on zero initialization on $\mathbf{v}$. We see a numerically convergent result, and the inner product between learned and true directions starts to grow from 0. As the directions guide the magnitude to grow, there is an extra stage for the directions to become roughly accurate. The choice of this initialization is necessary and subtle. The figure on the right is for small initialization 10^{-3} , where the entries outside support get significant magnitudes, and the algorithm fails.

E.2 AUTOENCODER WITH GROUPING LAYER

The grouping layers have been used in grouped CNN and grouped attention mechanisms (Wu et al., 2021; Xie et al., 2017; Lee et al., 2018), which usually leads to parameter efficiency and better accuracy. To demonstrate the practical value of such grouping layers, we conduct the following experiment about learning good representations on MNIST.

(Jing et al., 2020) proposed implicit rank-minimizing autoencoder (IRMAE), which is a deterministic autoencoder with implicit regularization. The idea is to apply more linear layers between encoder and decoder to penalize the rank of latent representation. A graphical illustration of the architecture is shown in Figure 8, where we explicitly show the last convolution layer and the linear layers in the latent space, which are absorbed into the last layer of the encoder in practice. This design is related to the power parametrization (Schwarz et al., 2021) trick to promote sparsity/low-rankness. One major advantage is that IRMAE produces a more interpretable latent representation, and the linear interpolation in the latent space gives a natural transition between two images.

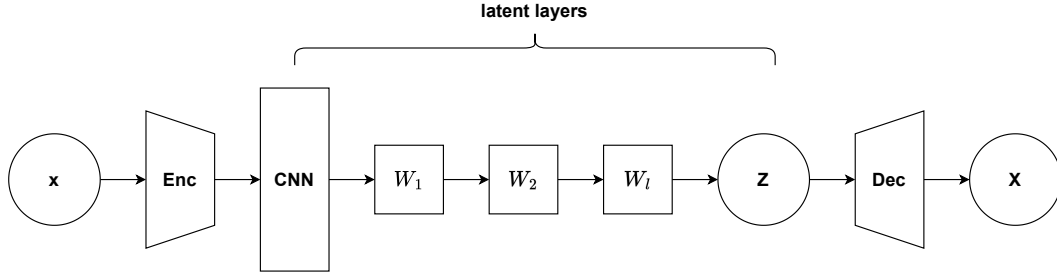


Figure 8: Implicit rank-minimizing autoencoder.

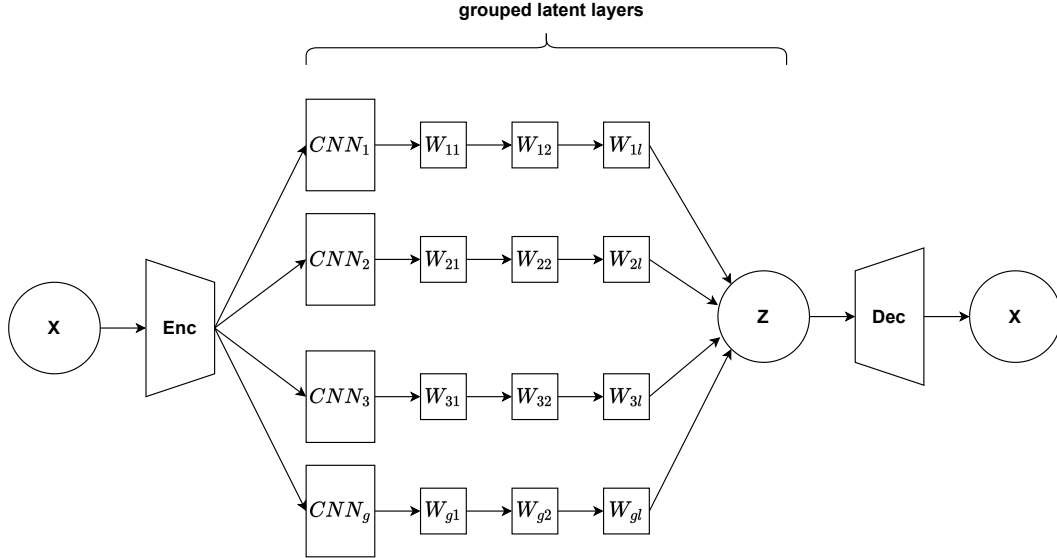


Figure 9: Implicit rank-minimizing autoencoder with grouping layers.

Inspired by our DGLNN, we design a CNN analog of it, which we call grouped autoencoder (GAE). The architecture is shown in Figure 9. The channels feed into the last convolutional layer of encoder is separable into g groups. The linear layers (power-parametrization) are applied within each group. Grouping channels of convolutional layers is a common practice to improve the parameter efficiency. With these grouping and power layers in the latent space, we expect it learns a better latent representation as IRMAE does.

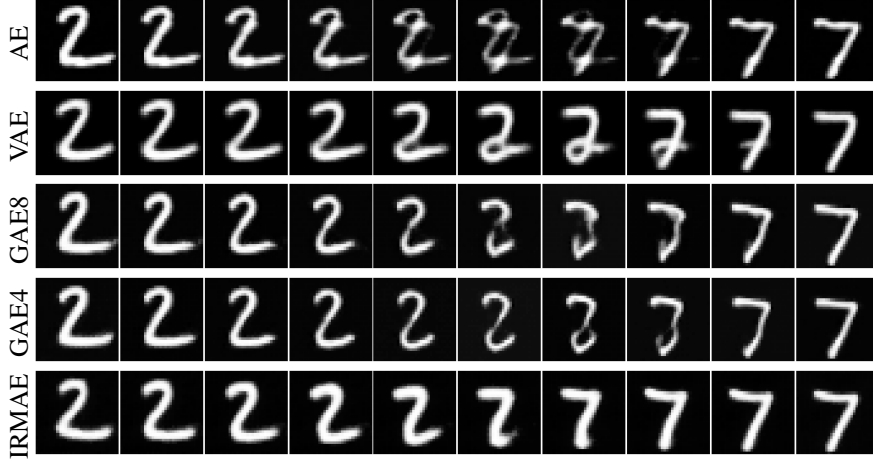


Figure 10: Linear interpolations between data points on the MNIST dataset. GAE4/8 stands for grouped autoencoder with 4/8 groups.

The linear interpolations between data points in the latent space are shown in Figure 10. We compare the grouped autoencoder (GAE) with autoencoder (AE), variational autoencoder (VAE) and implicit rank-minimizing autoencoder (IRMAE). We see that GAE outperforms AE and VAE, and gives comparable results with IRMAE. However, GAE achieves a better parameter efficiency as shown in Table 2.

	# of params
IRMAE	786K
GAE4	196K
GAE8	98K

Table 2: Number of parameters of hidden layers in latent space.

E.3 EXPERIMENTS WITH GAUSSIAN MEASUREMENTS

Besides the numerical results shown in Section 5, we conduct the following experiments with sampling each entry of $\mathbf{X}$ from a standard normal distribution.

The effectiveness. We follow the same setting with that Figure 3 except changing Rademacher random variables to Gaussian random variables. The convergence of Algorithm 2 is shown in Figure 11. We see that the recovered entries, group magnitudes and directions successfully converge to the true ones.

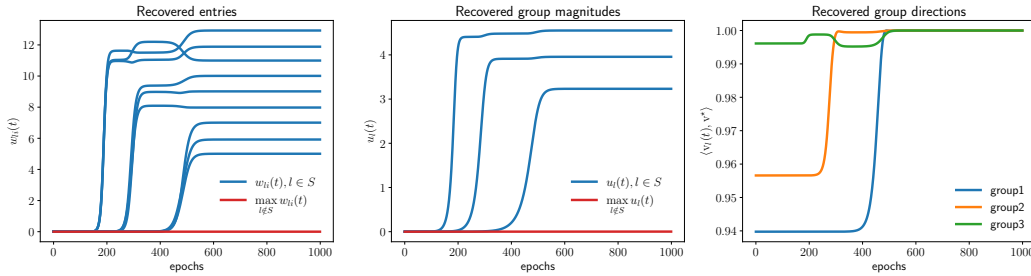


Figure 11: Convergence of algorithm 2 with Gaussian measurements

Comparisons with explicit regularization methods. We compare Algorithm 2 with proximal gradient descent implemented in (Carmichael et al., 2021) and primal-dual procedure (Molinari et al.,

2021). Each entry of $\mathbf{X}$ is sampled from a standard Gaussian distribution. We set $n = 150$ and $p = 300$, and the number of non-zero entries is 10, divided into 3 groups with size 4. We vary the variance in the noise to achieve different signal-to-noise ratios (SNR). The experiment is repeated 30 times at each noise level. The average and standard deviation of the estimation error are depicted in Figure 12. Our algorithm is consistently better than explicit regularization methods, whereas the primal-dual procedure has a comparable performance when SNR is large.

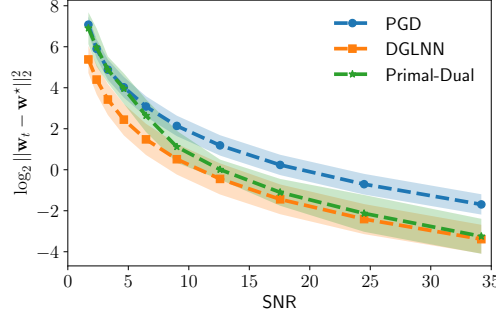


Figure 12: Comparisons with proximal gradient descent and iterative regularization.

To further discover the potential applications of our findings, we use a gene expression dataset from the Microarray experiments of mammalian eye tissue samples (Scheetz et al., 2006). The dataset consists of 120 samples with 100 predictors that are expanded from 20 genes using 5 basis B-splines, as described in (Yang & Zou, 2015). The goal is to predict the gene expression level of TRIM32, which causes Bardet-Biedl syndrome. We randomly split the data equally, and use the validation dataset for hyperparameter tuning and early stopping. We compare our approach with the commonly used proximal gradient descent and a primal-dual approach. The result is shown in Table 3. Our approach achieves the best performance among these three methods.

Test error	PGD	Primal-Dual	Our approach
MSE	0.03096	0.02868	0.02477

Table 3: Comparisons of MSE (mean squared error) on test set.