
Robust Training under Label Noise by Over-parameterization

Sheng Liu¹ Zhihui Zhu² Qing Qu³ Chong You⁴

Abstract

Recently, over-parameterized deep networks, with increasingly more network parameters than training samples, have dominated the performances of modern machine learning. However, when the training data is corrupted, it has been well-known that over-parameterized networks tend to overfit and do not generalize. In this work, we propose a principled approach for robust training of over-parameterized deep networks in classification tasks where a proportion of training labels are corrupted. The main idea is yet very simple: label noise is sparse and incoherent with the network learned from clean data, so we model the noise and learn to separate it from the data. Specifically, we model the label noise via another *sparse over-parameterization* term, and exploit implicit algorithmic regularizations to recover and separate the underlying corruptions. Remarkably, when trained using such a simple method in practice, we demonstrate state-of-the-art test accuracy against label noise on a variety of real datasets. Furthermore, our experimental results are corroborated by theory on simplified linear models, showing that exact separation between sparse noise and low-rank data can be achieved under incoherent conditions. The work opens many interesting directions for improving over-parameterized models by using sparse over-parameterization and implicit regularization. Code is available at <https://github.com/shengliu66/SOP>.

1. Introduction

One of the most important factors for the success of deep models is their large model size and high expressive power, which enable them to learn complicated input-output rela-

tions. As such, over-parametrized deep networks or large models, with more parameters than the size of training data, have dominated the performance in computer vision, natural language processing, and so on. The adoption of large models is justified by the recent discovery that deep models exhibit a “double descent” (Belkin et al., 2019) and “unimodal variance” (Yang et al., 2020) generalization behavior, where their performance continues to improve beyond the interpolation point, extending the classical learning theory of bias-variance trade-off. While there are infinitely many global solutions that *overfit* to training data, the choice of optimization algorithm imposes certain *implicit* regularization (Neyshabur et al., 2014) so that over-parameterized models converge to those that are generalizable.

Nonetheless, the success of over-parameterization of deep networks critically depends on the availability of *clean* training data, while overfitting inevitably occurs when training data is corrupted. Consider the task of image classification with a training dataset $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$, with $\mathbf{x}_i$ being an input image and $\mathbf{y}_i$ being the corresponding one-hot label. With an over-parameterized deep network $f(\cdot; \theta)$, model training is achieved by solving an optimization problem with respect to (w.r.t.) the network parameter θ as follows:

$$\min_{\theta} L(\theta) = \frac{1}{N} \sum_{i=1}^N \ell(f(\mathbf{x}_i; \theta), \mathbf{y}_i), \quad (1)$$

where $\ell(\cdot, \cdot)$ is a loss function that measures the distance between network prediction $f(\mathbf{x}_i; \theta)$ and the label $\mathbf{y}_i$. If a proportion of the images in the training set is *mislabelled* (Song et al., 2020), it is well-known that the network will be optimized to zero training error hence produce $f(\mathbf{x}_i; \theta) \approx \mathbf{y}_i$ for all $i \in \{1, \dots, N\}$, even for $\mathbf{y}_i$'s that are incorrect (Zhang et al., 2021a). Overfitting to wrong labels inevitably leads to poor generalization performance (see Fig. 1).

In this paper, we introduce a principled method to address the challenges of overfitting over-parameterized deep networks in the presence of training data corruptions. We focus on the task of classification trained with noisy label, a ubiquitous problem in practice due to the extreme complexity of data annotation even for experienced domain experts (Frénay & Verleysen, 2013). Our idea leverages the property that the label noise is *sparse*, namely only a fraction of the labels are corrupted and the rest are intact. Principled methods for dealing with sparse corruption have a rich

¹Center for Data Science, New York University ²Electrical and Computer Engineering, University of Denver ³Department of EECS, University of Michigan ⁴Google Research, New York City. Correspondence to: Chong You <cyou@google.com>.

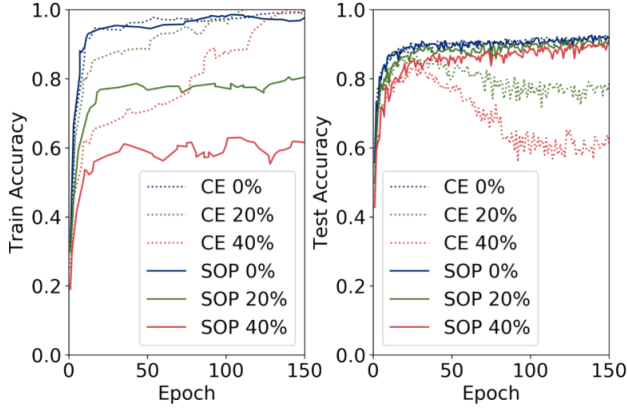


Figure 1. Sparse over-parameterization prevents overfitting to label noise. Training and test accuracy of a PreActResNet18 network trained with a standard cross entropy (CE) loss (dashed lines) and our Sparse Over-parameterization (SOP) (solid lines) for image classification on the CIFAR-10 dataset with 0%, 20%, and 40% of the labels flipped at random. SOP prevents overfitting to the wrong training labels, obtaining near 100%, 80%, 60% training accuracy respectively, therefore achieves better generalization on the test set without an accuracy drop at the end of training.

history, which can be retraced back to compressed sensing (Candes & Tao, 2005), robust subspace recovery (Candès et al., 2011; Wright et al., 2008), and even earlier (Claerbout & Muir, 1973). Such methods are based on using a robust loss function, such as the ℓ_1 norm which is less sensitive to large outlying entries. While it is tempting to use sparse modeling for the label noise problem by setting the loss $\ell()$ in (1) as the ℓ_1 loss, such an approach cannot solve the overfitting issue since all global solutions are still given by those that satisfy $f(x_i; \theta) \approx y_i$ for all $i \in \{1, \dots, N\}$. Hence, handling sparse corruptions with over-parameterized models requires the development of techniques beyond the classical ℓ_1 loss for sparse modeling.

Overview of our method and contribution. To handle sparse corruption with over-parameterized models, our idea is simply to use an extra variable s_i to model the unknown label noise s_{*i} , which is the difference between the observed label y_i and the corresponding clean label. Hence, the goal is to minimize the discrepancy between $f(x_i; \theta) + s_i$ and y_i . Inspired by a line of recent work (Vaskevicius et al., 2019; Zhao et al., 2019; You et al., 2020), we enforce sparsity of s_i by the over-parameterization $s_i = u_i \odot u_i - v_i \odot v_i$ and optimize the following training loss

$$\min_{\theta, \{u_i, v_i\}_{i=1}^N} L(\theta, \{u_i, v_i\}_{i=1}^N), \quad (2)$$

where

$$L(\theta, \{u_i, v_i\}) = \frac{1}{N} \sum_{i=1}^N \ell(f(x_i; \theta) + u_i \odot u_i - v_i \odot v_i, y_i),$$

with $\odot$ denoting an entry-wise Hadamard product. We term our method *Sparse Over-Parameterization* (SOP).

At the first glance, our SOP approach is seemingly problem-

atic, because adding more learnable parameters $\{u_i, v_i\}_{i=1}^N$ to an over-parameterized network $f(\cdot, \theta)$ would aggravate rather than alleviate the overfitting issue. Indeed, a global solution to (2) is given by $u_i \equiv v_i \equiv \mathbf{0}$ and $f(x_i, \theta) \equiv y_i$ for all $i \in \{1, \dots, N\}$ where the network overfits to noisy labels. Here, we leverage the choice of a particular training algorithm to enforce an *implicit bias* towards producing the desired solutions. Technically, we run gradient descent on the objective in (2) starting from a small initialization for $\{u_i, v_i\}_{i=1}^N$:

$$\begin{aligned} \theta &\leftarrow \theta - \tau \cdot \frac{\partial L(\theta, \{u_i, v_i\})}{\partial \theta}, \\ u_i &\leftarrow u_i - \alpha \tau \cdot \frac{\partial L(\theta, \{u_i, v_i\})}{\partial u_i}, \quad i = 1, \dots, N, \\ v_i &\leftarrow v_i - \alpha \tau \cdot \frac{\partial L(\theta, \{u_i, v_i\})}{\partial v_i}, \quad i = 1, \dots, N, \end{aligned} \quad (3)$$

where $\alpha > 0$ is the ratio of learning rates for different training variables. Such a simple algorithm enables our method of SOP to train a deep image classification networks without overfitting to wrong labels and obtain better generalization performance (see Fig. 1). A more comprehensive empirical study with a variety of datasets is presented in Section 2.

To rigorously justify our method, we theoretically investigate our method based upon a simplified over-parameterized linear model with sparse corruptions. As justified by a line of recent work (Jacot et al., 2018; Chizat et al., 2019), over-parameterized linear models capture similar phenomena because they well approximate over-parameterized deep networks in a linearized regime around the initial points. Under sparse corruption and certain low-rank assumptions on the data, we show that the gradient descent (3) with an α below a certain threshold recovers the underlying model parameters with sparse corruptions. Our result is obtained by explicitly characterizing the implicit regularization for the term $u_i \odot u_i - v_i \odot v_i$. In particular, we explicitly show that it leads to an ℓ_1 -norm regularization on the sparse corruption, hence connecting our method to classical ℓ_1 loss approaches for model robustness. For more details, we refer readers to Section 3.

In summary, our contributions are two-folds:

- *Method.* We proposed a simple yet practical SOP method that can effectively prevent overfitting for learning over-parameterized deep networks from corrupted training data, demonstrated on a variety of datasets.
- *Theory.* Under a simplified over-parameterized linear model, we rigorously justify our approach for exactly separating sparse corruption from the data.

Moreover, we believe the methodology we developed here could be far beyond the label noise setting, with the potential for dealing with more challenging scenarios of preventing overfitting in learning modern over-parametrized models of an ever-increasing size.

2. Robust Classification with Label Noise

In this section, we show how our SOP method plays out on image classification problems with the noisy label. In particular, we discuss extra implementation details of our method, followed by experimental demonstrations on a variety of datasets with synthetic and real label noise.

2.1. Implementation Details of SOP

We train an over-parameterized deep neural network $f(\cdot; \theta)$ from the noisy training data $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$ using the method briefly discussed in Section 1. Specifically, we train the network $f(\cdot; \theta)$ using the objective (2) with stochastic gradient descent (SGD) (i.e. a batch version of (3)).

Notice that there is additional prior information on label noise $\mathbf{s}_{*i}$ associated with a sample $\{\mathbf{x}_i, \mathbf{y}_i\}$, namely, the positive and negative entries of $\mathbf{s}_{*i}$ must correspond to nonzero entry and zero entries of $\mathbf{y}_i$, respectively. Moreover, all entries of $\mathbf{s}_{*i}$ must lie in the range of $[-1, 1]$. To leverage such information, we optimize a variant of (2) given by

$$\min_{\theta, \{\mathbf{u}_i, \mathbf{v}_i\}_{i=1}^N} \frac{1}{N} \sum_{i=1}^N \ell(f(\mathbf{x}_i; \theta) + \mathbf{s}_i, \mathbf{y}_i), \quad (4)$$

$$\text{s.t. } \mathbf{s}_i \doteq \mathbf{u}_i \odot \mathbf{u}_i \odot \mathbf{y}_i - \mathbf{v}_i \odot \mathbf{v}_i \odot (1 - \mathbf{y}_i), \text{ and} \quad (5)$$

$$\mathbf{u}_i \in [-1, 1]^K, \quad \mathbf{v}_i \in [-1, 1]^K, \quad (6)$$

where K is the number of classes. In above, constraints on $\mathbf{u}_i, \mathbf{v}_i$ are realized by performing a projection step after each gradient descent update.

Choice of the loss function $\ell(\cdot, \cdot)$ in (4). The most commonly used loss function for classification tasks is the cross-entropy loss $\ell_{\text{CE}}(\cdot, \cdot)$ (Krizhevsky et al., 2012). Because the $\ell_{\text{CE}}(\cdot, \cdot)$ loss requires a probability distribution as an input, we define a mapping

$$\phi(\mathbf{w}) \doteq \frac{\max\{\mathbf{w}, \epsilon \mathbf{1}\}}{\|\max\{\mathbf{w}, \epsilon \mathbf{1}\}\|_1}, \quad (7)$$

and set the loss $\ell(\cdot, \cdot)$ in (4) to be

$$L_{\text{CE}}(\theta, \mathbf{u}_i, \mathbf{v}_i; \mathbf{x}_i, \mathbf{y}_i) \doteq \ell_{\text{CE}}(\phi(f(\mathbf{x}_i; \theta) + \mathbf{s}_i), \mathbf{y}_i). \quad (8)$$

On the other hand, the cross-entropy loss cannot be used to optimize the variables $\{\mathbf{v}_i\}$ (see Section A.1 for an explanation). Hence, we use the mean squared error loss ℓ_{MSE} and set the loss in (4) to be

$$L_{\text{MSE}}(\theta, \mathbf{u}_i, \mathbf{v}_i; \mathbf{x}_i, \mathbf{y}_i) \doteq \ell_{\text{MSE}}(f(\mathbf{x}_i; \theta) + \mathbf{s}_i, \mathbf{y}_i), \quad (9)$$

when optimizing $\{\mathbf{v}_i\}$ ¹. We summarize our training method in Algorithm 1.

¹We also project $f(\mathbf{x}_i; \theta)$ to a one-hot vector when using MSE loss which is empirically found to accelerate convergence of $\{\mathbf{v}_i\}$.

Algorithm 1 Image classification under label noise by the method of Sparse Over-Parameterization (SOP).

- 1: **Input:** Training data $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$, network backbone $f(\cdot, \theta)$, variables $\{\mathbf{u}_i, \mathbf{v}_i\}_{i=1}^N$, number of epochs T , learning rate τ , learning rate ratio α_u, α_v , batch size β
- 2: **Initialization:** Draw entries of $\mathbf{u}_i, \mathbf{v}_i$ from i.i.d. Gaussian distribution with zero-mean and s.t.d. $1e - 8$
- 3: **for** each $t \in \{1, \dots, T\}$ **do**
- 4: # Train network $f(\cdot, \theta)$ with SGD
- 5: **for** each $b \in \{1, \dots, N/\beta\}$ **do**
- 6: Sample a batch $\mathcal{B} \subseteq \{1, \dots, N\}$ with $|\mathcal{B}| = \beta$
- 7: Set $\theta \leftarrow \theta - \tau \cdot \sum_{i \in \mathcal{B}} \frac{\partial L_{\text{CE}}(\theta, \mathbf{u}_i, \mathbf{v}_i; \mathbf{x}_i, \mathbf{y}_i)}{\partial \theta}$
- 8: **end for**
- 9: # Update $\{\mathbf{u}_i, \mathbf{v}_i\}$
- 10: **for** each $i \in \{1, \dots, N\}$ **do**
- 11: Set $\mathbf{u}_i \leftarrow \mathcal{P}_{[-1, 1]} \left(\mathbf{u}_i - \alpha_u \tau \frac{\partial L_{\text{CE}}(\theta, \mathbf{u}_i, \mathbf{v}_i; \mathbf{x}_i, \mathbf{y}_i)}{\partial \mathbf{u}_i} \right)$
- 12: Set $\mathbf{v}_i \leftarrow \mathcal{P}_{[-1, 1]} \left(\mathbf{v}_i - \alpha_v \tau \frac{\partial L_{\text{MSE}}(\theta, \mathbf{u}_i, \mathbf{v}_i; \mathbf{x}_i, \mathbf{y}_i)}{\partial \mathbf{v}_i} \right)$
- 13: **end for**
- 14: **end for**
- 15: **Output:** Network parameters θ and $\{\mathbf{u}_i, \mathbf{v}_i\}_{i=1}^N$

2.2. Experiments

We experimentally demonstrate the effectiveness of our proposed SOP method on datasets with both synthetic (i.e., CIFAR-10 and CIFAR-100) and realistic (i.e., CIFAR-N, Clothing-1M, and WebVision) label noise. In addition to the SOP described in Algorithm 1, we also implement an improved version, termed SOP+, which incorporates two commonly used regularization techniques in the literature of label noise, namely the consistency regularization and the class-balance regularization. We explain SOP+ in more detail in Appendix A.3.

Dataset descriptions. We use datasets with synthetic label noise generated from CIFAR-10 and CIFAR-100 datasets (Krizhevsky et al., 2009). Each dataset contains 50k training images and 10k test images, all with clean labels, where each image is of size 32×32 . Following previous works (Han et al., 2018; Liu et al., 2020; Xia et al., 2020), we generate symmetric label noise by uniformly flipping labels for a percentage of the training set for all classes, as well as asymmetric label noise by flipping labels for particular pairs of classes. For datasets with realistic label noise, we test on CIFAR-10N/CIFAR-100N (Wei et al., 2021b) which contains a re-annotation of CIFAR-10/CIFAR-100 with human workers. Specifically, each image in CIFAR-10N contains three submitted labels (i.e., **Random 1, 2, 3**) which are further combined to have an **Aggregate** and a **Worst** label. Each image in CIFAR-100N contains a single submitted label for the fine classes. We also test on Clothing-1M (Xiao et al., 2015) which is a large-scale dataset with images clawed from online shopping websites and labels

Table 1. Test accuracy with synthetic label noise on CIFAR-10 and CIFAR-100 with {20%, 40%, 60%, 80%} percent of labels for training data randomly flipped uniformly to another class. All methods use ResNet34 as the architecture. Mean and standard deviation over 5 independent runs are reported.

METHODS	CIFAR-10				CIFAR-100			
	20%	40%	60%	80%	20%	40%	60%	80%
CE	86.32 $\pm$ 0.18	82.65 $\pm$ 0.16	76.15 $\pm$ 0.32	59.28 $\pm$ 0.97	51.43 $\pm$ 0.58	45.23 $\pm$ 0.53	36.31 $\pm$ 0.39	20.23 $\pm$ 0.82
FORWARD	87.99 $\pm$ 0.36	83.25 $\pm$ 0.38	74.96 $\pm$ 0.65	54.64 $\pm$ 0.44	39.19 $\pm$ 2.61	31.05 $\pm$ 1.44	19.12 $\pm$ 1.95	8.99 $\pm$ 0.58
GCE	89.83 $\pm$ 0.20	87.13 $\pm$ 0.22	82.54 $\pm$ 0.23	64.07 $\pm$ 1.38	66.81 $\pm$ 0.42	61.77 $\pm$ 0.24	53.16 $\pm$ 0.78	29.16 $\pm$ 0.74
SL	89.83 $\pm$ 0.32	87.13 $\pm$ 0.26	82.81 $\pm$ 0.61	68.12 $\pm$ 0.81	70.38 $\pm$ 0.13	62.27 $\pm$ 0.22	54.82 $\pm$ 0.57	25.91 $\pm$ 0.44
ELR	91.16 $\pm$ 0.08	89.15 $\pm$ 0.17	86.12 $\pm$ 0.49	73.86 $\pm$ 0.61	74.21 $\pm$ 0.22	68.28 $\pm$ 0.31	59.28 $\pm$ 0.67	29.78 $\pm$ 0.56
SOP (OURS)	93.18 $\pm$ 0.57	90.09 $\pm$ 0.27	86.76 $\pm$ 0.22	68.32 $\pm$ 0.77	74.67 $\pm$ 0.30	70.12 $\pm$ 0.57	60.26 $\pm$ 0.41	30.20 $\pm$ 0.63

Table 2. Comparison with the state-of-the-art methods that use two network ensembles and semi-supervised learning on CIFAR-10 and CIFAR-100 under symmetric (with 20%, 50%, 80%) and asymmetric (with 40%) label noise. All methods use ResNet34 as the architecture.

METHODS	CIFAR-10				CIFAR-100			
	SYMMETRIC			ASYM	SYMMETRIC			ASYM
	20%	50%	80%	40%	20%	50%	80%	40%
CE	87.2	80.7	65.8	82.2	58.1	47.1	23.8	43.3
MixUP	93.5	87.9	72.3	-	69.9	57.3	33.6	-
DIVIDEMIX	96.1	94.6	93.2	93.4	77.1	74.6	60.2	72.1
ELR+	95.8	94.8	93.3	93.0	77.7	73.8	60.8	77.5
SOP+ (OURS)	96.3	95.5	94.0	93.8	78.8	75.9	63.3	78.0

Table 3. Test accuracy with realistic label noise on Clothing1M and WebVision. We use a pre-trained ResNet50 for Clothing1M and an InceptionResNetV2 for WebVision dataset. The results of the comparing methods are taken from their respective papers.

METHODS	CLOTHING1M	WEBVISION	ILSVRC12
CE	69.1	-	-
FORWARD	69.8	61.1	57.3
CO-TEACHING	69.2	63.6	61.5
ELR	72.9	76.2	68.7
CORES ²	73.2	-	-
SOP (OURS)	73.5	76.6	69.1

generated based on surrounding texts. Clothing-1M contains 1 million training images, 15k validation images, and 10k test images with clean labels. Finally, we also test on the mini WebVision dataset (Li et al., 2017) which contains the top 50 classes from the Google image subset of WebVision (approximately 66 thousand images). Models trained on mini WebVision are evaluated on both WebVision and ImageNet ILSVRC12 validation set. Details on the label noise for these datasets is provided in Section A.2.

Network structures & hyperparameters. We implement our method with PyTorch v1.7. For each dataset, the choices of network architectures and hyperparameters for SOP are as follows. Additional details, as well as hyper-parameters for both SOP and SOP+, can be found in Appendix A.4.

- *CIFAR-10/100 and CIFAR-10N/100N.* We follow (Liu et al., 2020) to use ResNet-34 and PreActResNet18 architectures trained with SGD using a 0.9 momentum. The initial learning rate is 0.02 decayed with a factor of 10

at the 40th and 80th epochs for CIFAR-10/CIFAR-10N and at the 80th and 120th epochs for CIFAR-100/CIFAR-100N, respectively. Weight decay for network parameters θ is set to 5×10^{-4} . No weight decay is used for parameters $\{u_i, v_i\}_{i=1}^N$.

- *Clothing-1M.* We follow the previous work (Liu et al., 2020) to use a ResNet-50 (He et al., 2016) pre-trained on ImageNet (Krizhevsky et al., 2012). The network is trained with batch size 64 and an initial learning rate 0.001, which is reduced by a factor of 10 after 5th epoch (10 training epochs in total). Optimization is performed using SGD with a momentum 0.9. Weight decay is 0.001 for parameters θ and is zero for parameters $\{u_i, v_i\}_{i=1}^N$.
- *Mini Webvision.* We use InceptionResNetV2 as the backbone architecture. All other optimization details are the same as for CIFAR-10, except that we use weight decay 0.0005 and batch size 32.

Experimental results. We compare with methods based on estimation of the transition matrix (Forward (Patrini et al., 2017)), design of loss functions (GCE (Zhang & Sabuncu, 2018) and SL (Wang et al., 2019)), training two networks (Co-teaching (Han et al., 2018) and DivideMix (Li et al., 2020a)), and label noise correction (ELR (Liu et al., 2020) and CORES² (Cheng et al., 2021)).

Table 1 reports the performance of our method on synthetically generated symmetric label noise using CIFAR-10 and CIFAR-100. To compare with state-of-the-art methods, we also report the performance of SOP+ which contains additional regularization on both symmetric and asymmetric label noise and report the results in Table 2. It can be observed that our method is robust to a fairly large amount of label noise, and compares favorably to existing techniques.

We further demonstrate that our method can effectively handle datasets with realistic label noise by reporting its performance on Clothing1M & WebVision (see Table 3) and CIFAR-N (see Table 4) datasets. We can observe a performance gain over all comparing methods.

Finally, we compare the training time (on a single Nvidia V100 GPU) of our method to the baseline methods in Table 5. We observe that our algorithm SOP/SOP+ achieves the fastest speed across all baselines.

Table 4. **Test accuracy with realistic label noise on CIFAR-100N.** Mean and standard deviation over 5 independent runs are reported. The results of the baseline methods are taken from (Wei et al., 2021b) which all use ResNet34 as the architecture. For SOP+, we use PreActResNet18.

METHODS	CIFAR-100N						CIFAR-100N	
	CLEAN	RANDOM 1	RANDOM 2	RANDOM 3	AGGREGATE	WORST	CLEAN	NOISY
CE	92.92 $\pm$ 0.11	85.02 $\pm$ 0.65	86.46 $\pm$ 1.79	85.16 $\pm$ 0.61	87.77 $\pm$ 0.38	77.69 $\pm$ 1.55	76.70 $\pm$ 0.74	55.50 $\pm$ 0.66
FORWARD	93.02 $\pm$ 0.12	86.88 $\pm$ 0.50	86.14 $\pm$ 0.24	87.04 $\pm$ 0.35	88.24 $\pm$ 0.22	79.79 $\pm$ 0.46	76.18 $\pm$ 0.37	57.01 $\pm$ 1.03
CO-TEACHING	93.35 $\pm$ 0.14	90.33 $\pm$ 0.13	90.30 $\pm$ 0.17	90.15 $\pm$ 0.18	91.20 $\pm$ 0.13	83.83 $\pm$ 0.13	73.46 $\pm$ 0.09	60.37 $\pm$ 0.27
ELR+	95.39 $\pm$ 0.05	94.43 $\pm$ 0.41	94.20 $\pm$ 0.24	94.34 $\pm$ 0.22	94.83 $\pm$ 0.10	91.09 $\pm$ 1.60	78.57 $\pm$ 0.12	66.72 $\pm$ 0.07
CORES*	94.16 $\pm$ 0.11	94.45 $\pm$ 0.14	94.88 $\pm$ 0.31	94.74 $\pm$ 0.03	95.25 $\pm$ 0.09	91.66 $\pm$ 0.09	73.87 $\pm$ 0.16	55.72 $\pm$ 0.42
SOP+(OURS)	96.38 $\pm$ 0.31	95.28 $\pm$ 0.13	95.31 $\pm$ 0.10	95.39 $\pm$ 0.11	95.61 $\pm$ 0.13	93.24 $\pm$ 0.21	78.91 $\pm$ 0.43	67.81 $\pm$ 0.23

Table 5. **Comparison of total training time** in hours on CIFAR-10 with 50% symmetric label noise.

CE	CO-TEACHING+	DIVIDEMIX	ELR+	SOP	SOP+
0.9H	4.4H	5.4H	2.3H	1.0H	2.1H

3. Theoretical Insights with Simplified Models

This section provides theoretical insights into our SOP method by studying structured data recovery with sparse corruption in the context of over-parameterized *linear* models. We will start with model simplification, followed by our main theoretical results and experimental verification.

3.1. Problem Setup & Main Result

Given a highly *overparameterized* network $f(\cdot; \theta)$, recent work (Jacot et al., 2018; Kalimeris et al., 2019) suggests that the parameter $\theta \in \mathbb{R}^p$ may not change much from its initialization θ_0 before obtaining zero training error. Hence, a nonlinear network $f(\cdot; \theta) : \mathbb{R}^n \mapsto \mathbb{R}$ can be well approximated by its first-order Taylor expansion:

$$f(x; \theta) \approx f(x; \theta_0) + \langle \nabla_{\theta} f(x; \theta_0), \theta - \theta_0 \rangle, \quad (10)$$

where we consider $f(\cdot; \theta)$ as a scalar function for simplicity. Since the bias term $f(x; \theta_0) - \langle \nabla_{\theta} f(x; \theta_0), \theta_0 \rangle$ is constant w.r.t. θ , for simplicity we may further assume that

$$f(x; \theta) \approx \langle \nabla_{\theta} f(x; \theta_0), \theta \rangle. \quad (11)$$

Thus, for a dataset $\{x_i\}_{i=1}^N$ of N points, collectively

$$\begin{bmatrix} f(x_1; \theta) \\ \vdots \\ f(x_N; \theta) \end{bmatrix} \approx \begin{bmatrix} \nabla_{\theta} f(x_1; \theta_0) \\ \vdots \\ \nabla_{\theta} f(x_N; \theta_0) \end{bmatrix} \cdot \theta = J \cdot \theta, \quad (12)$$

where $J \in \mathbb{R}^{N \times p}$ is a Jacobian matrix. This observation motivates us to consider the following problem setup.

Problem setup. Based upon the above linearization, we assume that our corrupted observation $y \in \mathbb{R}^N$ (e.g., noisy labels) is generated by

$$y = J \cdot \theta_* + s_*, \quad (13)$$

where $\theta_* \in \mathbb{R}^p$ is the underlying groundtruth parameter, and the noise $s_* \in \mathbb{R}^N$ is *sparse* so that only a subset of observation (e.g., labels) is corrupted. Given J and y

generated from (13), our goal is to recover both θ_* and s_* .

However, as we are considering the problem in an over-parameterized regime with $p > N$, the underdetermined system (13) implies that there are *infinite* solutions for θ_* even if s_* is given. Nonetheless, recent work showed that the implicit bias of gradient descent for overparameterized linear models and deep networks tend to find minimum ℓ_2 -norm solutions (Zhang et al., 2021a). To make our problem more well-posed, motivated by these results, we would like to find an θ_* with minimum ℓ_2 -norm, namely,

$$\theta_* = \arg \min_{\theta} \|\theta\|_2^2 \quad \text{s.t.} \quad y = J\theta + s_*. \quad (14)$$

Analogous to (2), we will show that θ_* and s_* can be provably recovered by solving the problem

$$\min_{\theta, u, v} h(\theta, u, v) \doteq \frac{1}{2} \|J\theta + u \odot u - v \odot v - y\|_2^2, \quad (15)$$

using the gradient descent algorithm with learning rates τ and $\alpha\tau$ on θ and $\{u, v\}$, respectively:

$$\begin{aligned} \theta_{k+1} &= \theta_k - \tau \cdot J^\top r_k, \\ u_{k+1} &= u_k - 2\alpha\tau \cdot u_k \odot r_k, \\ v_{k+1} &= v_k + 2\alpha\tau \cdot v_k \odot r_k, \end{aligned} \quad (16)$$

where $r_k \doteq J\theta_k + u_k \odot u_k - v_k \odot v_k - y$. Based on these, our result can be summarized as follows.

Theorem 3.1 (Main result, informal). *Suppose J is rank- r and μ -incoherent defined in Section 3.3, and s_* is k -sparse. If $k^2 r < N/(4\mu)$, with $\tau \rightarrow 0$ and a proper choice of α depending on (J, θ_*, k) , the gradient dynamics of (16) converges to the ground truth solution (θ_*, s_*) in (13) starting from a small initialization of (θ, u, v) .*

We state our result at a high level with more technical details in Section 3.2 and Section 3.3. The overall idea of the proof can be sketched through the following two steps.

- First, although the problem (15) is *nonconvex*, in Section 3.2 we show that it has benign global landscape, and that the gradient descent (16) converges to particular global solutions that are the same as solutions to a *convex* problem with explicit regularizations on θ and s .
- Building upon above results, in Section 3.3 we complete our analysis by showing that θ_* and s_* can be exactly

recovered by the convex problem with a small enough value for α .

Throughout the analysis, we corroborate our findings with numerical simulations.

3.2. Landscapes & Implicit Sparse Regularization

Benign global landscape. We start by characterizing the nonconvex landscape of (15), showing the following result.

Proposition 3.2. *Any critical point of (15) is either a global minimizer, or it is a strict saddle (Ge et al., 2015) with its Hessian having at least one negative eigenvalue.*

For a strict saddle function, recent work (Lee et al., 2016) showed that gradient descent with random initialization almost surely escapes saddle points and converges to a local minimizer. Thus, Proposition 3.2 ensures that the algorithm in (16) almost surely converges to a global solution of (15).

However, because there are infinite many global solutions for the overparameterized model (15) and not all global solutions are of equal quality, convergence to a global solution alone is not sufficient for us to establish the correctness of our method. Nonetheless, as we will show in the following, the particular choice of the algorithm in (16) enables it to converge to a particular *regularized* global solution.

Implicit sparse regularization. To understand which solution the algorithm (16) converges to, we study its *gradient flow* counterpart by taking the stepsize $\tau \rightarrow 0$ in (16). Thus, the dynamics of such a gradient flow is governed by the following differential equations

$$\begin{aligned}\dot{\theta}_t(\gamma, \alpha) &= -\mathbf{J}^\top \mathbf{r}_t(\gamma, \alpha), \\ \dot{\mathbf{u}}_t(\gamma, \alpha) &= -2\alpha \cdot \mathbf{u}_t(\gamma, \alpha) \odot \mathbf{r}_t(\gamma, \alpha), \\ \dot{\mathbf{v}}_t(\gamma, \alpha) &= 2\alpha \cdot \mathbf{v}_t(\gamma, \alpha) \odot \mathbf{r}_t(\gamma, \alpha),\end{aligned}\quad (17)$$

where we define

$$\mathbf{r}_t(\gamma, \alpha) = \mathbf{J}\theta_t(\gamma, \alpha) + \mathbf{u}_t(\gamma, \alpha) \odot \mathbf{u}_t(\gamma, \alpha) - \mathbf{v}_t(\gamma, \alpha) \odot \mathbf{v}_t(\gamma, \alpha) - \mathbf{y}. \quad (18)$$

Here, we assume that θ , $\mathbf{u}$, and $\mathbf{v}$ are initialized at

$$\theta_0(\gamma, \alpha) = \mathbf{0}, \mathbf{u}_0(\gamma, \alpha) = \gamma \mathbf{1}, \mathbf{v}_0(\gamma, \alpha) = \gamma \mathbf{1}, \quad (19)$$

with some small $\gamma > 0$. Solving the differential equations in (17) gives the gradient flow

$$\begin{aligned}\theta_t(\gamma, \alpha) &= \mathbf{J}^\top \boldsymbol{\nu}_t(\gamma, \alpha), \\ \mathbf{u}_t(\gamma, \alpha) &= \gamma \exp(2\alpha \boldsymbol{\nu}_t(\gamma, \alpha)), \\ \mathbf{v}_t(\gamma, \alpha) &= \gamma \exp(-2\alpha \boldsymbol{\nu}_t(\gamma, \alpha)),\end{aligned}\quad (20)$$

where we define

$$\boldsymbol{\nu}_t(\gamma, \alpha) \doteq - \int_0^t \mathbf{r}_\tau(\gamma, \alpha) d\tau. \quad (21)$$

The following result shows that the solution that the gradient flow $(\theta_t(\gamma, \alpha), \mathbf{u}_t(\gamma, \alpha), \mathbf{v}_t(\gamma, \alpha))$ in (20) converges to at $t \rightarrow \infty$ is a global solution to (15) that is *regularized* with a

particular of (γ, α) .

Proposition 3.3. *Consider the gradient flow in (20) with the initialization in (19).*

- **(Global convergence)** *For any (γ, α) , if the limit*

$$\begin{aligned}(\theta_\infty(\gamma, \alpha), \mathbf{u}_\infty(\gamma, \alpha), \mathbf{v}_\infty(\gamma, \alpha)) \\ \doteq \lim_{t \rightarrow \infty} (\theta_t(\gamma, \alpha), \mathbf{u}_t(\gamma, \alpha), \mathbf{v}_t(\gamma, \alpha))\end{aligned}\quad (22)$$

exists, then it is a global solution to (15).

- **(Implicit regularization)** *Fix any $\lambda > 0$ and let α be a function of γ as*

$$\alpha(\gamma) = -\frac{\log \gamma}{2\lambda}. \quad (23)$$

If the limit

$$\begin{aligned}(\hat{\theta}, \hat{\mathbf{u}}, \hat{\mathbf{v}}) \doteq \lim_{\gamma \rightarrow 0} (\theta_\infty(\gamma, \alpha(\gamma)), \\ \mathbf{u}_\infty(\gamma, \alpha(\gamma)), \mathbf{v}_\infty(\gamma, \alpha(\gamma)))\end{aligned}\quad (24)$$

exists, then $(\hat{\theta}, \hat{\mathbf{u}}, \hat{\mathbf{v}})$ is a global solution to (15). In particular, let

$$\hat{\mathbf{s}} \doteq \hat{\mathbf{u}} \odot \hat{\mathbf{u}} - \hat{\mathbf{v}} \odot \hat{\mathbf{v}}, \quad (25)$$

then $(\hat{\theta}, \hat{\mathbf{s}})$ is an optimal solution to the convex program

$$\min_{\theta, \mathbf{s}} \frac{1}{2} \|\theta\|_2^2 + \lambda \|\mathbf{s}\|_1, \quad \text{s.t. } \mathbf{y} = \mathbf{J}\theta + \mathbf{s}. \quad (26)$$

As we observe from the above result, because $(\hat{\theta}, \hat{\mathbf{s}})$ that the gradient flow (17) converges to is also an optimal solution of (26), it implies that $(\hat{\theta}, \hat{\mathbf{s}})$ is regularized. In particular, the ℓ_1 -norm regularization on $\mathbf{s}$ comes as a result of implicit regularization on overparameterization $\mathbf{s} = \mathbf{u} \odot \mathbf{u} - \mathbf{v} \odot \mathbf{v}$, leading to a sparse solution on $\mathbf{s}$ as we desired. On the other hand, the ℓ_2 regularization on θ leads to the desired minimum ℓ_2 -norm solution as we discussed in (14). Thus, the only question remains is whether the ground truth $(\theta_*, \mathbf{s}_*)$ in (13) can be identified through solving the convex problem (26), which we will discuss in the following Section 3.3.

Numerical verification. While Proposition 3.3 is proved for gradient flow with both learning rate $\tau \rightarrow 0$ and initialization scale $\gamma \rightarrow 0$, we numerically show that such a result also holds non-asymptotically with finitely small τ and γ .

Given a tuple (N, p, r, k) of model parameters, we generate simulation data $(\mathbf{J}, \theta_*, \mathbf{s}_*, \mathbf{y})$ as follows. The matrix $\mathbf{J} \in \mathbb{R}^{N \times p}$ is generated by multiplying two randomly generated matrices of shape $N \times r$ and $r \times p$, respectively with entries drawn i.i.d. from a standard Gaussian distribution. The sparse vector $\mathbf{s}_* \in \mathbb{R}^N$ is generated by randomly choosing k entries to be i.i.d. standard Gaussian, with the rest of entries zero. Then, we generate a vector $\theta \in \mathbb{R}^p$ with all entries drawn i.i.d. from a standard Gaussian distribution, and let $\mathbf{y} = \mathbf{J}\theta + \mathbf{s}_*$. Finally, we set θ_* as the minimum ℓ_2 -norm solution according to (14).

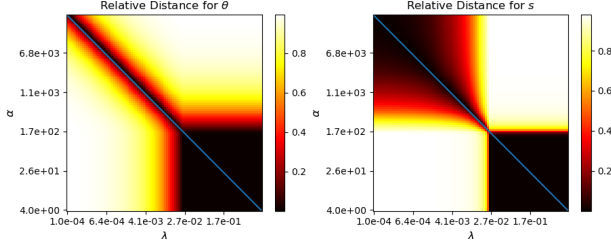


Figure 2. The gradient descent in (16) and the convex problem in (26) produce the same solutions with $\alpha = -\frac{\log \gamma}{2\lambda}$. For fixed data $(\mathbf{J}, \mathbf{y})$, left figure shows the relative difference $\frac{\|\theta_\alpha - \theta_\lambda\|_2}{\max\{\|\theta_\alpha\|_2, \|\theta_\lambda\|_2\}}$ between the solution θ_α to (16) with varying values of α (in y-axis) and the solution θ_λ computed from (26) with varying values of λ (in x-axis). Likewise, right figure shows the relative difference for $\mathbf{s}$. Blue line shows the curve $\alpha = -\frac{\log \gamma}{2\lambda}$ where γ is fixed to $\exp(-8)$ in all experiments.

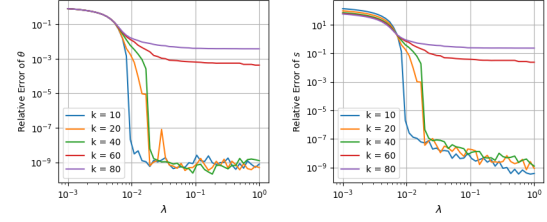
In this experiment, we choose and fix $(N, p, r, k) = (20, 40, 3, 3)$ for the data generation described above. With a varying learning rate $\alpha \in [4, 4000]$, we compute $(\theta_\alpha, \mathbf{s}_\alpha)$ as the solution provided by gradient descent in (16) with an initialization by (19) with $\gamma = e^{-8}$. With a varying regularization $\lambda \in [0.0001, 1]$ in (26), we compute $(\theta_\lambda, \mathbf{s}_\lambda)$ as the solution provided by the convex problem in (26) with weight parameter λ , using the ECOS solver (Domahidi et al., 2013) provided in CVXPY (Diamond & Boyd, 2016). Figure 2 provides a visualization of the relative difference $\rho = \frac{\|\theta_\alpha - \theta_\lambda\|_2}{\max\{\|\theta_\alpha\|_2, \|\theta_\lambda\|_2\}}$ between θ_α and θ_λ (and likewise for $\mathbf{s}$), across all pairs of (α, λ) . We can observe that as long as (α, λ) satisfies the relationship in (23), the relative difference ρ is small for θ , which is also true for $\mathbf{s}$. On the other hand, the relative differences can be large if (23) is not satisfied, corroborating Proposition 3.3.

3.3. Exact Recovery under Incoherence Conditions

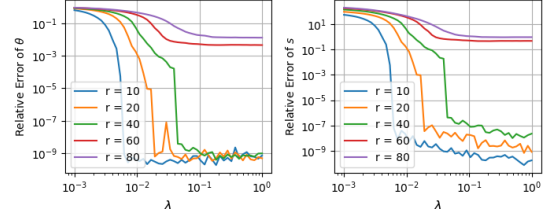
Given the overparameterized model (13) with $\mathbf{y} \in \mathbb{R}^N$, $\theta \in \mathbb{R}^p$, and $p \gg N$, there is not enough information from $\mathbf{y}$ to recover θ_* and $\mathbf{s}_*$ even with the prior information that $\mathbf{s}_*$ is sparse – any given vector $\mathbf{y} \in \mathbb{R}^p$ can be decomposed as a summation of an arbitrary sparse vector $\mathbf{s}$ and a vector θ cooked up from the column space of $\mathbf{J}$ as long as $\mathbf{J}$ has full row rank.

For the solution θ_* and $\mathbf{s}_*$ to be identified, first, we assume that $\mathbf{J}$ is low-rank, motivated by empirical observation in practical deep neural network f_θ that the Jacobian matrix $\mathbf{J}$ of f_θ is approximately low-rank (Oymak et al., 2019).² However, the low-rank condition of $\mathbf{J}$ alone does not guarantee identifiability, because it cannot address the separability

²The low-rank assumption simplifies our analysis but at the cost that our model is not able to overfit to any corrupted labels as a deep neural network. We leave the study under a realistic approximate low-rank assumption to future work.



(a) Varying k with fixed $r = 20$.



(b) Varying r with fixed $k = 20$.

Figure 3. Effect of model parameter λ for exact recovery by (26). The y-axis is the relative error of θ (left) and $\mathbf{s}$ (right) defined as $\frac{\|\theta - \theta_*\|_2}{\|\theta_*\|_2}$ and $\frac{\|\mathbf{s} - \mathbf{s}_*\|_2}{\|\mathbf{s}_*\|_2}$, respectively, where $(\theta, \mathbf{s})$ is the solution to (26). The curves are averages over 10 independent trials.

between $\mathbf{J}\theta_*$ and $\mathbf{s}_*$ – following a similar argument as that in (Candès et al., 2011), if any column of $\mathbf{J}$ has a single nonzero entry, then any $\mathbf{s}_*$ that is supported on the same entry cannot be recovered without ambiguities. Hence, we further assume that the column space of $\mathbf{J}$ and the standard basis $[e_1, \dots, e_N] \doteq \text{diag}\{1, \dots, 1\} \in \mathbb{R}^{N \times N}$ are incoherent, defined as follows.

Definition 3.4 (Candès et al. (2011)). Let $\mathbf{J} = \mathbf{U}\Sigma\mathbf{V}^\top \in \mathbb{R}^{N \times p}$ be the compact SVD of $\mathbf{J}$ and r be the rank of $\mathbf{J}$. The coherence of $\mathbf{J}$ (w.r.t. the standard basis) is defined as

$$\mu(\mathbf{J}) = \frac{N}{r} \max_{1 \leq i \leq N} \|\mathbf{U}^\top e_i\|_2^2. \quad (27)$$

It should be noted that the low-rank and incoherence assumptions are common for matrix recovery (Davenport & Romberg, 2016; Chi et al., 2019). Based upon the above assumptions on $\mathbf{J}$ and $\mathbf{s}_*$, we show the following.

Proposition 3.5. Let r be the rank of $\mathbf{J}$ and k be the number of nonzero entries of $\mathbf{s}_*$. If we have

$$k^2 r < \frac{N}{4\mu(\mathbf{J})}, \quad (28)$$

then the solution to (26) is $(\theta_*, \mathbf{s}_*)$ for any $\lambda > \lambda_0$, where $\lambda_0 > 0$ is a scalar depending on $(\mathbf{J}, \theta_*, k)$.

Thus, combining this result with Proposition 3.3, the gradient flow in (17) with initialization (19) converges to $(\theta_*, \mathbf{s}_*)$ when the choice of learning rate ratio α in (17) is smaller than a certain threshold, justifying our claim in Theorem 3.1.

Numerical verification. To corroborate Proposition 3.5, we numerically solve (26) under varying conditions of λ , r , and k . The simulated data $(\mathbf{J}, \theta_*, \mathbf{s}_*, \mathbf{y})$ is generated

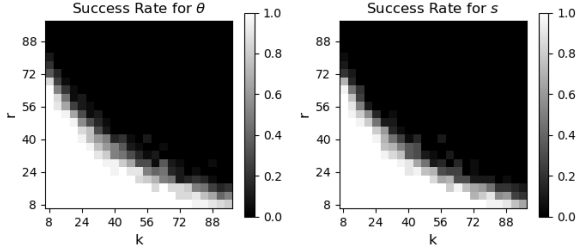


Figure 4. Phase transition for solving (26) over 20 trials, with fixed $\lambda = 0.1$ and varying k, r . Recovery is declared success if $\frac{\|\theta - \theta_*\|_2}{\|\theta_*\|_2} < 0.001$ (left) and $\frac{\|s - s_*\|_2}{\|s_*\|_2} < 0.001$ (right).

the same way as the experimental part in Section 3.2 with $N = 100$ and $p = 150$, and for an obtained solution (θ, s) via solving (26), we measure the relative recovery error $\epsilon_\theta = \frac{\|\theta - \theta_*\|_2}{\|\theta_*\|_2}$ and $\epsilon_s = \frac{\|s - s_*\|_2}{\|s_*\|_2}$.

- *Effects of the parameter λ .* Here, we consider the recovery with varying $\lambda \in [0.0001, 1]$. First, we fix $r = 20$ and vary $k \in \{10, 20, 40, 60, 80\}$, showing the relative recovery errors ϵ_θ and ϵ_s in Figure 3(a). Second, we fix $k = 20$ and vary $r \in \{10, 20, 40, 60, 80\}$, showing the results in Figure 3(b). The results show a clear phase transition that correct recovery is obtained only when λ is greater than a particular threshold λ_0 . Moreover, λ_0 varies depending on k and r , consistent with Proposition 3.5.
- *Relationships between the rank r and sparsity k .* Here, we fix $\lambda = 0.1$ and plot the phase transition with respect to r and k . For each (r, k) , the simulation is repeated for 20 random instances, and for each instance we declare the recovery to be successful if $\epsilon_\theta < 0.001$ and $\epsilon_s < 0.001$. As shown in Figure 4, the phase transition is consistent with Proposition 3.5 that successful recovery is achieved only when both k and r are small.

4. Related Work and Discussion

4.1. Prior Arts on Implicit Regularization

Since overparameterized deep neural networks do not overfit (in the absence of data corruption) even without any explicit regularization (Zhang et al., 2021a), it is argued that there are implicit regularizations of learning algorithms that enable the models to converge to desired solutions. Under the assumption of linear or deep linear models, many works characterized the mathematics of such implicit bias via explicit regularizations (Soudry et al., 2018; Gunasekar et al., 2018; Li et al., 2018; Oymak & Soltanolkotabi, 2019; Arora et al., 2019; Razin & Cohen, 2020; Li et al., 2020c; Ji et al., 2020; Stöger & Soltanolkotabi, 2021; Jacot et al., 2021). Among those, the closest related to ours include (Vaskevicius et al., 2019; Zhao et al., 2019; Woodworth et al., 2020; Li et al., 2021a; Chou et al., 2021), which studied the implicit *sparse* regularization induced by a term of the form $u \odot u - v \odot v$.

While all the above works aim to understand implicit regularization by studying linear models, the practical benefits of such studies are unclear. Our work provides an inspiring result showing that principled design with implicit regularization leads to robust learning of over-parameterized models. In particular, our model in (2) is motivated by existing studies on the implicit sparse regularization, but adds such a regularization to an (already) implicitly regularized model for handling sparse corruptions. In other words, two forms of implicit regularization are involved in our model which poses new problems in the design of the optimization algorithm and in mathematical analysis. To the best of our knowledge, the only prior works that use implicit sparse regularization for robust learning are (You et al., 2020; Ma & Fattahi, 2021; Ding et al., 2021) which studied the robust recovery of low-rank matrices and images. Among them, our work extends (You et al., 2020) to the problem of image classification with label noise, demonstrates its effectiveness, and provides dedicated theoretical analyses. Additionally, methods in (Ma & Fattahi, 2021; Ding et al., 2021) require a particular learning rate schedule that may not be compatible with commonly used schedules such as cosine annealing (Loshchilov & Hutter, 2016) in image classification.

4.2. Relationship to Existing Work on Label Noise

Deep neural networks are over-parameterized hence prone to *overfitting* to the label noises. While many popular regularization techniques for alleviating overfitting, such as label smoothing (Szegedy et al., 2016; Lukasik et al., 2020; Wei et al., 2021a) and *mixup* (Zhang et al., 2018), are useful for mitigating the impact of label noise, they do not completely solve the problem due to a lack of precise noise modeling. In the following, we discuss three of the most popular line of work dedicated to the label noise problem; we refer the reader to the survey papers (Algan & Ulusoy, 2021; Song et al., 2020; Wei et al., 2021b) for a comprehensive review.

Loss design. Robust loss function, such as the ℓ_1 loss (Ghosh et al., 2017), is one of the most popular approaches to the label noise problem which has many recent extensions (Zhang & Sabuncu, 2018; Wang et al., 2019; Amid et al., 2019; Ma et al., 2020; Yu et al., 2020; Wei & Liu, 2021). The method is based on reducing the loss associated with large outlying entries, hence the impact of label noise. A similar idea is also explored in gradient clipping (Menon et al., 2019) and loss reweighting (Liu & Tao, 2015; Wang et al., 2017; Chang et al., 2017; Zhang et al., 2021b; Zetterqvist et al., 2021) methods. While robust loss enables the model to learn faster from correct labels, *the global solution still overfits to corrupted labels with over-parameterized models*.

Label transition probability. Another popular line of work for label noise is based on the assumption that the noisy label is drawn from a probability distribution conditioned on

the true label. Here, the main task is to estimate the underlying transition probabilities. The early work (Chen & Gupta, 2015; Goldberger & Ben-Reuven, 2017) encapsulates the transition probabilities as a noise adaptation layer that is stacked on top of a classification network and trained jointly in an end-to-end fashion. Recent work (Patrini et al., 2017) uses separated procedures to estimate the transition probabilities, the success of which requires either the availability of a clean validation data (Hendrycks et al., 2018) or additional data assumptions (Xia et al., 2019; Zhu et al., 2021; Li et al., 2021b; Zhang et al., 2021c). Even if the underlying transition probabilities can be correctly recovered, *overfitting is only prevented asymptotically, requiring sufficiently many samples of corrupted labels for each input* (Patrini et al., 2017), which is not practical.

Label correction. In contrast to the above methods, our method completely avoids overfitting even with finite training samples. This is achieved by the over-parameterization term $u \odot u - v \odot v$ in (2) which recovers the clean labels. Hence, our method is related to techniques based on noisy label detection and refurbishment. Nonetheless, existing techniques are based on heuristic argument about different behaviors of clean and corrupted samples in the training process, such as properties of learned representations (Kim et al., 2021; Ma et al., 2018; Jiang et al., 2020), prediction consistency (Reed et al., 2014; Song et al., 2019a), learning speed (Li et al., 2020b; Liu et al., 2020; 2021a), margin (Lin & Bradic, 2021), confidence (Cheng et al., 2021). They often need to be combined with engineering tricks such as moving average (Huang et al., 2020; Liu et al., 2020) and burning-in (Zheng et al., 2020) to make them work well. Finally, the work (Hu et al., 2019) introduces a variable to estimate the label noise in a way similar to (2). However, the variable is not over-parameterized to induce sparsity, and their method does not have competitive performance.

4.3. Sparsity in Deep Learning

Our method is broadly related to existing efforts on introducing sparsity into deep learning (Hoeffler et al., 2021), but is notably different in both the objective of introducing sparsity, the origin of sparsity, and how sparsity is enforced. First, previous exploration of sparsity primarily aims to improve training and inference efficiency with large-scale models, while our paper focuses on robust training under label noise. Second, previous introduction of sparsity is often motivated by its presence in biological brains, but there is still a lack of clean understanding of how sparsity helps with learning. In contrast, sparsity in our method has the clear mathematical meaning that the percentage of corrupted labels is small. Finally, while pruning (Liu et al., 2021b; Chen et al., 2021) is a dominant approach for obtaining sparsity, our method leverages the implicit bias of gradient descent associated with a particular sparse over-parameterization.

4.4. Limitations and Future Directions

Choice of optimization algorithms. Our SOP method is based on introducing more parameters to an already over-parameterized model, hence relies critically on the choice of the optimization algorithm to induce the desired implicit regularization. For *vanilla* gradient descent, our analysis in Section 3 shows that it has the desired implicit regularization by design. In practical deep network training, it is more common to use the *stochastic* gradient descent *with momentum*. While not theoretically justified, experiments in Section 2 show that our method works with such practical variants. This may not come as a surprise, because existing studies already show that stochastic gradient descent (Nacson et al., 2019) and momentum acceleration (Wang et al., 2021) have the same implicit bias as the vanilla gradient descent under certain models. We leave the extension of such results to our method as future works.

Modeling of label noise. Our method is based on the assumption that the label noise matrix $S_* = [s_{*1}, \dots, s_{*N}]$, where s_{*i} is the difference between the observed label y_i and the underlying true label, is a *sparse* matrix. We made no additional assumption on the sparsity pattern of S_* , other than the non-negative and non-positive constraints discussed in (4). In practice, it is usually the case that certain pairs of classes are more similar hence more easily confusing with each other than other pairs. As a result, certain blocks of S_* tend to have more non-zero entries than the others. When there is a prior on which blocks may have more non-zero entries, our method may be adapted by using a *weighted* sparse regularization for the corresponding blocks. When there is no such prior, our method may be adapted by using a *group* sparse regularization (Neyshabur et al., 2014; Tibshirani, 2021).

Acknowledgements

SL and QQ were partially supported by NSF grant DMS 2009752. SL was partially supported by NSF NRT-HDR Award 1922658 and Alzheimer’s Association grant AARG-NTF-21-848627. QQ also acknowledge support from NSF CAREER 2143904, NSF CCF 2212066, and ONR N00014-22-1-2529. ZZ acknowledges support from NSF grants CCF 2008460 and CCF 2106881. Part of this work was done when CY was at University of California, Berkeley and was supported by Tsinghua-Berkeley Shenzhen Institute Research Fund. The authors acknowledge helpful discussion with Ryan Chan from Johns Hopkins University.

References

- Algan, G. and Ulusoy, I. Image classification with deep learning in the presence of noisy labels: A survey. *Knowledge-Based Systems*, 215:106771, 2021.

- Amid, E., Warmuth, M. K., Anil, R., and Koren, T. Robust bi-tempered logistic loss based on bregman divergences. *Advances in Neural Information Processing Systems*, 32: 15013–15022, 2019.
- Arora, S., Cohen, N., Hu, W., and Luo, Y. Implicit regularization in deep matrix factorization. In *Advances in Neural Information Processing Systems*, pp. 7411–7422, 2019.
- Belkin, M., Hsu, D., Ma, S., and Mandal, S. Reconciling modern machine-learning practice and the classical bias–variance trade-off. *Proceedings of the National Academy of Sciences*, 116(32):15849–15854, 2019.
- Berthelot, D., Carlini, N., Goodfellow, I., Papernot, N., Oliver, A., and Raffel, C. A. Mixmatch: A holistic approach to semi-supervised learning. In Wallach, H., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- Candes, E. J. and Tao, T. Decoding by linear programming. *IEEE transactions on information theory*, 51(12):4203–4215, 2005.
- Candès, E. J., Li, X., Ma, Y., and Wright, J. Robust principal component analysis? *Journal of the ACM (JACM)*, 58(3): 1–37, 2011.
- Chang, H.-S., Learned-Miller, E., and McCallum, A. Active bias: Training more accurate neural networks by emphasizing high variance samples. *Advances in Neural Information Processing Systems*, 30:1002–1012, 2017.
- Chen, T., Zhang, Z., Balachandra, S., Ma, H., Wang, Z., Wang, Z., et al. Sparsity winning twice: Better robust generalization from more efficient training. In *International Conference on Learning Representations*, 2021.
- Chen, X. and Gupta, A. Webly supervised learning of convolutional networks. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 1431–1439, 2015.
- Cheng, H., Zhu, Z., Li, X., Gong, Y., Sun, X., and Liu, Y. Learning with instance-dependent label noise: A sample sieve approach. In *International Conference on Learning Representations*, 2021.
- Chi, Y., Lu, Y. M., and Chen, Y. Nonconvex optimization meets low-rank matrix factorization: An overview. *IEEE Transactions on Signal Processing*, 67(20):5239–5269, 2019.
- Chizat, L., Oyallon, E., and Bach, F. On lazy training in differentiable programming. *Advances in Neural Information Processing Systems*, 32:2937–2947, 2019.
- Chou, H.-H., Maly, J., and Rauhut, H. More is less: Inducing sparsity via overparameterization. *arXiv preprint arXiv:2112.11027*, 2021.
- Claerbout, J. F. and Muir, F. Robust modeling with erratic data. *Geophysics*, 38(5):826–844, 1973.
- Cohen, A., Dahmen, W., and DeVore, R. Compressed sensing and best k -term approximation. *Journal of the American mathematical society*, 22(1):211–231, 2009.
- Davenport, M. A. and Romberg, J. An overview of low-rank matrix recovery from incomplete observations. *IEEE Journal of Selected Topics in Signal Processing*, 10(4): 608–622, 2016.
- Diamond, S. and Boyd, S. CVXPY: A Python-embedded modeling language for convex optimization. *Journal of Machine Learning Research*, 17(83):1–5, 2016.
- Ding, L., Jiang, L., Chen, Y., Qu, Q., and Zhu, Z. Rank over-specified robust matrix recovery: Subgradient method and exact recovery. *Advances in Neural Information Processing Systems*, 34, 2021.
- Domahidi, A., Chu, E., and Boyd, S. ECOS: An SOCP solver for embedded systems. In *European Control Conference (ECC)*, pp. 3071–3076, 2013.
- Frénay, B. and Verleysen, M. Classification in the presence of label noise: a survey. *IEEE transactions on neural networks and learning systems*, 25(5):845–869, 2013.
- Ge, R., Huang, F., Jin, C., and Yuan, Y. Escaping from saddle points—online stochastic gradient for tensor decomposition. In *Conference on learning theory*, pp. 797–842. PMLR, 2015.
- Ghosh, A., Kumar, H., and Sastry, P. Robust loss functions under label noise for deep neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 31, 2017.
- Goldberger, J. and Ben-Reuven, E. Training deep neural networks using a noise adaptation layer. 2017.
- Gunasekar, S., Woodworth, B., Bhojanapalli, S., Neyshabur, B., and Srebro, N. Implicit regularization in matrix factorization. In *2018 Information Theory and Applications Workshop (ITA)*, pp. 1–10. IEEE, 2018.
- Han, B., Yao, Q., Yu, X., Niu, G., Xu, M., Hu, W., Tsang, I., and Sugiyama, M. Co-teaching: Robust training of deep neural networks with extremely noisy labels. In Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018.

- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Hendrycks, D., Mazeika, M., Wilson, D., and Gimpel, K. Using trusted data to train deep networks on labels corrupted by severe noise. *Advances in Neural Information Processing Systems*, 31:10456–10465, 2018.
- Hoefler, T., Alistarh, D., Ben-Nun, T., Dryden, N., and Peste, A. Sparsity in deep learning: Pruning and growth for efficient inference and training in neural networks. *Journal of Machine Learning Research*, 22(241):1–124, 2021.
- Hu, W., Li, Z., and Yu, D. Simple and effective regularization methods for training on noisily labeled data with generalization guarantee. In *International Conference on Learning Representations*, 2019.
- Huang, L., Zhang, C., and Zhang, H. Self-adaptive training: beyond empirical risk minimization. *arXiv preprint arXiv:2002.10319*, 2020.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. *arXiv preprint arXiv:1806.07572*, 2018.
- Jacot, A., Ged, F., Gabriel, F., Simsek, B., and Hongler, C. Deep linear networks dynamics: Low-rank biases induced by initialization scale and l2 regularization. *arXiv preprint arXiv:2106.15933*, 2021.
- Ji, Z., Dudík, M., Schapire, R. E., and Telgarsky, M. Gradient descent follows the regularization path for general losses. In *Conference on Learning Theory*, pp. 2109–2136. PMLR, 2020.
- Jiang, L., Huang, D., Liu, M., and Yang, W. Beyond synthetic noise: Deep learning on controlled noisy labels. In *International Conference on Machine Learning*, pp. 4804–4815. PMLR, 2020.
- Kalimeris, D., Kaplun, G., Nakkiran, P., Edelman, B., Yang, T., Barak, B., and Zhang, H. Sgd on neural networks learns functions of increasing complexity. *Advances in Neural Information Processing Systems*, 32:3496–3506, 2019.
- Kim, T., Ko, J., Choi, J., Yun, S.-Y., et al. Fine samples for learning with noisy labels. *Advances in Neural Information Processing Systems*, 34, 2021.
- Krizhevsky, A., Hinton, G., et al. Learning multiple layers of features from tiny images. 2009.
- Krizhevsky, A., Sutskever, I., and Hinton, G. E. Imagenet classification with deep convolutional neural networks. *Advances in neural information processing systems*, 25: 1097–1105, 2012.
- Lee, J. D., Simchowitz, M., Jordan, M. I., and Recht, B. Gradient descent only converges to minimizers. In *Conference on learning theory*, pp. 1246–1257. PMLR, 2016.
- Li, J., Wong, Y., Zhao, Q., and Kankanhalli, M. S. Learning to learn from noisy labeled data. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 5051–5059, 2019.
- Li, J., Socher, R., and Hoi, S. C. Dividemix: Learning with noisy labels as semi-supervised learning. *arXiv preprint arXiv:2002.07394*, 2020a.
- Li, J., Nguyen, T., Hegde, C., and Wong, K. W. Implicit sparse regularization: The impact of depth and early stopping. *Advances in Neural Information Processing Systems*, 34, 2021a.
- Li, M., Soltanolkotabi, M., and Oymak, S. Gradient descent with early stopping is provably robust to label noise for overparameterized neural networks. In *International conference on artificial intelligence and statistics*, pp. 4313–4324. PMLR, 2020b.
- Li, W., Wang, L., Li, W., Agustsson, E., and Van Gool, L. Webvision database: Visual learning and understanding from web data. *arXiv preprint arXiv:1708.02862*, 2017.
- Li, X., Liu, T., Han, B., Niu, G., and Sugiyama, M. Provably end-to-end label-noise learning without anchor points. *arXiv preprint arXiv:2102.02400*, 2021b.
- Li, Y., Ma, T., and Zhang, H. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Conference On Learning Theory*, pp. 2–47, 2018.
- Li, Z., Luo, Y., and Lyu, K. Towards resolving the implicit bias of gradient descent for matrix factorization: Greedy low-rank learning. In *International Conference on Learning Representations*, 2020c.
- Lin, J. Z. and Bradic, J. Learning to combat noisy labels via classification margins. *arXiv preprint arXiv:2102.00751*, 2021.
- Liu, S., Niles-Weed, J., Razavian, N., and Fernandez-Granda, C. Early-learning regularization prevents memorization of noisy labels. *Advances in Neural Information Processing Systems*, 33, 2020.
- Liu, S., Liu, K., Zhu, W., Shen, Y., and Fernandez-Granda, C. Adaptive early-learning correction for segmentation from noisy annotations. *ArXiv*, abs/2110.03740, 2021a.

- Liu, S., Yin, L., Mocanu, D. C., and Pechenizkiy, M. Do we actually need dense over-parameterization? in-time over-parameterization in sparse training. In *International Conference on Machine Learning*, pp. 6989–7000. PMLR, 2021b.
- Liu, T. and Tao, D. Classification with noisy labels by importance reweighting. *IEEE Transactions on pattern analysis and machine intelligence*, 38(3):447–461, 2015.
- Loshchilov, I. and Hutter, F. Sgdr: Stochastic gradient descent with warm restarts. *arXiv preprint arXiv:1608.03983*, 2016.
- Lukasik, M., Bhojanapalli, S., Menon, A., and Kumar, S. Does label smoothing mitigate label noise? In *International Conference on Machine Learning*, pp. 6448–6458. PMLR, 2020.
- Ma, J. and Fattahi, S. Implicit regularization of sub-gradient method in robust matrix recovery: Don’t be afraid of outliers. *arXiv preprint arXiv:2102.02969*, 2021.
- Ma, X., Wang, Y., Houle, M. E., Zhou, S., Erfani, S., Xia, S., Wijewickrema, S., and Bailey, J. Dimensionality-driven learning with noisy labels. In *International Conference on Machine Learning*, pp. 3355–3364. PMLR, 2018.
- Ma, X., Huang, H., Wang, Y., Romano, S., Erfani, S., and Bailey, J. Normalized loss functions for deep learning with noisy labels. In *International Conference on Machine Learning*, pp. 6543–6553. PMLR, 2020.
- Menon, A. K., Rawat, A. S., Reddi, S. J., and Kumar, S. Can gradient clipping mitigate label noise? In *International Conference on Learning Representations*, 2019.
- Nacson, M. S., Srebro, N., and Soudry, D. Stochastic gradient descent on separable data: Exact convergence with a fixed learning rate. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 3051–3059. PMLR, 2019.
- Neyshabur, B., Tomioka, R., and Srebro, N. In search of the real inductive bias: On the role of implicit regularization in deep learning. *arXiv preprint arXiv:1412.6614*, 2014.
- Oymak, S. and Soltanolkotabi, M. Overparameterized non-linear learning: Gradient descent takes the shortest path? In *International Conference on Machine Learning*, pp. 4951–4960, 2019.
- Oymak, S., Fabian, Z., Li, M., and Soltanolkotabi, M. Generalization guarantees for neural networks via harnessing the low-rank structure of the jacobian. *arXiv preprint arXiv:1906.05392*, 2019.
- Patrini, G., Rozza, A., Krishna Menon, A., Nock, R., and Qu, L. Making deep neural networks robust to label noise: A loss correction approach. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 1944–1952, 2017.
- Razin, N. and Cohen, N. Implicit regularization in deep learning may not be explainable by norms. *Advances in Neural Information Processing Systems*, 33, 2020.
- Reed, S., Lee, H., Anguelov, D., Szegedy, C., Erhan, D., and Rabinovich, A. Training deep neural networks on noisy labels with bootstrapping. *arXiv preprint arXiv:1412.6596*, 2014.
- Song, H., Kim, M., and Lee, J.-G. Selfie: Refurbishing unclean samples for robust deep learning. In *International Conference on Machine Learning*, pp. 5907–5915. PMLR, 2019a.
- Song, H., Kim, M., Park, D., and Lee, J.-G. Prestopping: How does early stopping help generalization against label noise? *ArXiv*, abs/1911.08059, 2019b.
- Song, H., Kim, M., Park, D., Shin, Y., and Lee, J.-G. Learning from noisy labels with deep neural networks: A survey. *arXiv preprint arXiv:2007.08199*, 2020.
- Soudry, D., Hoffer, E., Nacson, M. S., Gunasekar, S., and Srebro, N. The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research*, 19(1):2822–2878, 2018.
- Stöger, D. and Soltanolkotabi, M. Small random initialization is akin to spectral learning: Optimization and generalization guarantees for overparameterized low-rank matrix reconstruction. *Advances in Neural Information Processing Systems*, 34, 2021.
- Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J., and Wojna, Z. Rethinking the inception architecture for computer vision. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 2818–2826, 2016.
- Tanaka, D., Ikami, D., Yamasaki, T., and Aizawa, K. Joint optimization framework for learning with noisy labels. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 5552–5560, 2018.
- Tibshirani, R. J. Equivalences between sparse models and neural networks. 2021.
- Vaskevicius, T., Kanade, V., and Rebeschini, P. Implicit regularization for optimal sparse recovery. In *Advances in Neural Information Processing Systems*, pp. 2968–2979, 2019.

- Wang, B., Meng, Q., Zhang, H., Sun, R., Chen, W., and Ma, Z.-M. Momentum doesn't change the implicit bias. *arXiv preprint arXiv:2110.03891*, 2021.
- Wang, R., Liu, T., and Tao, D. Multiclass learning with partially corrupted labels. *IEEE transactions on neural networks and learning systems*, 29(6):2568–2580, 2017.
- Wang, Y., Ma, X., Chen, Z., Luo, Y., Yi, J., and Bailey, J. Symmetric cross entropy for robust learning with noisy labels. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 322–330, 2019.
- Wei, J. and Liu, Y. When optimizing f -divergence is robust with label noise. In *International Conference on Learning Representations*, 2021.
- Wei, J., Liu, H., Liu, T., Niu, G., and Liu, Y. Understanding (generalized) label smoothing when learning with noisy labels. *arXiv preprint arXiv:2106.04149*, 2021a.
- Wei, J., Zhu, Z., Cheng, H., Liu, T., Niu, G., and Liu, Y. Learning with noisy labels revisited: A study using real-world human annotations. *arXiv preprint arXiv:2110.12088*, 2021b.
- Woodworth, B., Gunasekar, S., Lee, J. D., Moroshko, E., Savarese, P., Golan, I., Soudry, D., and Srebro, N. Kernel and rich regimes in overparametrized models. In *Conference on Learning Theory*, pp. 3635–3673. PMLR, 2020.
- Wright, J., Yang, A. Y., Ganesh, A., Sastry, S. S., and Ma, Y. Robust face recognition via sparse representation. *IEEE transactions on pattern analysis and machine intelligence*, 31(2):210–227, 2008.
- Xia, X., Liu, T., Wang, N., Han, B., Gong, C., Niu, G., and Sugiyama, M. Are anchor points really indispensable in label-noise learning? *Advances in Neural Information Processing Systems*, 32:6838–6849, 2019.
- Xia, X., Liu, T., Han, B., Gong, C., Wang, N., Ge, Z., and Chang, Y. Robust early-learning: Hindering the memorization of noisy labels. In *International Conference on Learning Representations*, 2020.
- Xiao, T., Xia, T., Yang, Y., Huang, C., and Wang, X. Learning from massive noisy labeled data for image classification. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 2691–2699, 2015.
- Xie, Q., Dai, Z., Hovy, E. H., Luong, M.-T., and Le, Q. V. Unsupervised data augmentation for consistency training. *arXiv: Learning*, 2020.
- Yang, Z., Yu, Y., You, C., Steinhardt, J., and Ma, Y. Rethinking bias-variance trade-off for generalization of neural networks. In *International Conference on Machine Learning*, pp. 10767–10777. PMLR, 2020.
- You, C., Zhu, Z., Qu, Q., and Ma, Y. Robust recovery via implicit bias of discrepant learning rates for double over-parameterization. *Advances in Neural Information Processing Systems*, 33:17733–17744, 2020.
- Yu, Y., Chan, K. H. R., You, C., Song, C., and Ma, Y. Learning diverse and discriminative representations via the principle of maximal coding rate reduction. *Advances in Neural Information Processing Systems*, 33:9422–9434, 2020.
- Zetterqvist, O., Jörnsten, R., and Jonasson, J. Robust neural network classification via double regularization. *arXiv preprint arXiv:2112.08102*, 2021.
- Zhang, C., Bengio, S., Hardt, M., Recht, B., and Vinyals, O. Understanding deep learning (still) requires rethinking generalization. *Communications of the ACM*, 64(3):107–115, 2021a.
- Zhang, H., Cisse, M., Dauphin, Y. N., and Lopez-Paz, D. mixup: Beyond empirical risk minimization. In *International Conference on Learning Representations*, 2018.
- Zhang, H., Xing, X., and Liu, L. Dualgraph: A graph-based method for reasoning about label noise. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 9654–9663, 2021b.
- Zhang, Y., Niu, G., and Sugiyama, M. Learning noise transition matrix from only noisy labels via total variation regularization. *arXiv preprint arXiv:2102.02414*, 2021c.
- Zhang, Z. and Sabuncu, M. R. Generalized cross entropy loss for training deep neural networks with noisy labels. In *32nd Conference on Neural Information Processing Systems (NeurIPS)*, 2018.
- Zhao, P., Yang, Y., and He, Q.-C. Implicit regularization via hadamard product over-parametrization in high-dimensional linear regression. *arXiv preprint arXiv:1903.09367*, 2019.
- Zheng, S., Wu, P., Goswami, A., Goswami, M., Metaxas, D., and Chen, C. Error-bounded correction of noisy labels. In *International Conference on Machine Learning*, pp. 11447–11457. PMLR, 2020.
- Zhu, Z., Song, Y., and Liu, Y. Clusterability as an alternative to anchor points when learning with noisy labels. *arXiv preprint arXiv:2102.05291*, 2021.

Appendices

This appendix is organized as follows. In Section A we provide additional details for reproducing experimental results presented in Section 2. In Section B we provide proofs for the theoretical results presented in Section 3.

A. Training Details for Robust Classification with Label Noise

A.1. Choice of Loss Function

The cross-entropy loss in (8) cannot be used to optimize $\{v_i\}$ as we explain below. Consider a data point $\mathbf{x}$ with a one-hot label $\mathbf{y}$. With the CE loss in (8) rewritten below for convenience:

$$L_{\text{CE}}(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}; \mathbf{x}, \mathbf{y}) \doteq \ell_{\text{CE}}\left(\phi(f(\mathbf{x}, \boldsymbol{\theta}) + \mathbf{s}), \mathbf{y}\right), \text{ with } \mathbf{s} \doteq \mathbf{u} \odot \mathbf{u} \odot \mathbf{y} - \mathbf{v} \odot \mathbf{v} \odot (1 - \mathbf{y}), \quad (\text{A.1})$$

we may compute its gradient with respect to (w.r.t.) $\mathbf{v}$ as

$$\frac{\partial L_{\text{CE}}(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}; \mathbf{x}, \mathbf{y})}{\partial \mathbf{v}} = \frac{2\mathbf{v} \odot (1 - \mathbf{y})}{\mathbf{1}^\top (f(\mathbf{x}, \boldsymbol{\theta}) + \mathbf{s})}. \quad (\text{A.2})$$

This shows that the gradient w.r.t. different entries of $\mathbf{v}$ does not depend on the output $f(\mathbf{x}, \boldsymbol{\theta})$ of the model at all modulo the divider shared by all entries. Hence, $\mathbf{v}$ cannot correctly learn the label noise.

We now consider the MSE loss in (9) rewritten below for convenience:

$$L_{\text{MSE}}(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}; \mathbf{x}, \mathbf{y}) \doteq \ell_{\text{MSE}}\left(f(\mathbf{x}, \boldsymbol{\theta}) + \mathbf{s}, \mathbf{y}\right), \text{ with } \mathbf{s} \doteq \mathbf{u} \odot \mathbf{u} \odot \mathbf{y} - \mathbf{v} \odot \mathbf{v} \odot (1 - \mathbf{y}). \quad (\text{A.3})$$

The gradient w.r.t. $\mathbf{v}$ can be computed as

$$\frac{\partial L_{\text{MSE}}(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}; \mathbf{x}, \mathbf{y})}{\partial \mathbf{v}} = 4(f(\mathbf{x}, \boldsymbol{\theta}) + \mathbf{s} - \mathbf{y}) \odot \mathbf{v} \odot (1 - \mathbf{y}). \quad (\text{A.4})$$

Here the gradient w.r.t. different entries of $\mathbf{v}$ varies depending on how well the model prediction $f(\mathbf{x}, \boldsymbol{\theta}) + \mathbf{s}$ matches the given label $\mathbf{y}$ at the corresponding entry. Hence, when the model prediction deviates from the given label which may occur when the label is corrupted, $\mathbf{v}$ is able to learn the underlying corruption to the label.

A.2. Definition of Label Noise

In this paper, we consider two types of widely existed label noise, namely symmetric label noise and asymmetric label noise. For symmetric noise with noise level α , the labels are generated as follows:

$$\mathbf{y} = \begin{cases} \mathbf{y}^{GT} & \text{with the probability of } 1 - \alpha \\ \text{random one hot vector} & \text{with the probability of } \alpha. \end{cases}$$

We consider noise level $\alpha \in \{0.2, 0.4, 0.5, 0.6, 0.8\}$. For asymmetric noise, following (Patrini et al., 2017), we flip labels between TRUCK $\rightarrow$ AUTOMOBILE, BIRD $\rightarrow$ AIRPLANE, DEER $\rightarrow$ HORSE, and CAT $\leftrightarrow$ DOG. We randomly choose 40% training data with their labels to be flipped according to this asymmetric labeling rule. For real world datasets, Clothing1M has noise level estimated at around 38.5% (Song et al., 2019b), and for WebVision, the noise level is estimated to be at around 20% (Li et al., 2017).

A.3. Implementation Details of SOP+

We considered two separate regularization terms to further boost the results and stabilize training. We will describe the definitions and roles of them below:

Consistency regularizer $\mathcal{L}_C$. We use a regularizer $\mathcal{L}_C$ to encourage consistency of network prediction on a original image and the corresponding augmented image. Such a regularizer is commonly used in semi-supervised learning and label noise learning literature, see e.g., (Berthelot et al., 2019; Li et al., 2019). Specifically, the consistency regularizer $\mathcal{L}_C$ is defined as the Kullback-Leibler (KL)-divergence between the softmax predictions from the images with augmentations (described in Section A.4) and the softmax predictions for the corresponding images generated with Unsupervised Data Augmentation

Table A.1. Hyper-parameters for SOP on CIFAR-10/100, Clothing-1M and Webvision datasets.

	CIFAR-10		CIFAR-100		Clothing-1M	Webvision
architecture	ResNet34	PreAct PresNet18	ResNet34	PreAct PresNet18	ResNet-50 (pretrained)	InceptionResNetV2
batch size	128	128	128	128	64	32
learning rate (lr)	0.02	0.02	0.02	0.02	0.002	0.02
lr decay	40th & 80th	Cosine Annealing	40th & 80th	Cosine Annealing	5th	50th
weight decay (wd)	5×10^{-4}	5×10^{-4}	5×10^{-4}	5×10^{-4}	1×10^{-3}	5×10^{-4}
training epochs	120	300	150	300	10	100
training examples	45,000	50,000	45,000	50,000	1,000,000	66,000
lr for $\{u_i, v_i\}$	Sym: $\alpha_u = 10, \alpha_v = 10$ Asym: $\alpha_u = 10, \alpha_v = 100$	Sym: $\alpha_u = 10, \alpha_v = 10$ Asym: $\alpha_u = 10, \alpha_v = 100$	Sym: $\alpha_u = 1, \alpha_v = 10$ Asym: $\alpha_u = 1, \alpha_v = 100$	Sym: $\alpha_u = 1, \alpha_v = 10$ Asym: $\alpha_u = 1, \alpha_v = 100$	$\alpha_u = 0.1, \alpha_v = 1$	$\alpha_u = 0.1, \alpha_v = 1$
wd for $\{u_i, v_i\}$	0	0	0	0	0	0
init. std for $\{u_i, v_i\}$	10^{-8}	10^{-8}	10^{-8}	10^{-8}	10^{-8}	10^{-8}
λ_C	0.0	0.9	0.0	0.9	0.0	0.0
λ_B	0.0	0.1	0.0	0.1	0.0	0.0

(UDA) (Xie et al., 2020):

$$\mathcal{L}_c(\theta) = \frac{1}{N} \sum_{i=1}^N D_{KL}(f(x_i; \theta) \parallel f(\text{UDA}(x_i); \theta)).$$

Class-balance regularizer $\mathcal{L}_B$. We use a regularizer $\mathcal{L}_B$ to prevent the network from assigning all data points to the same class. Following (Tanaka et al., 2018), we use the prior information on the probability distribution p of class labels and minimize its distance in terms of KL-divergence to the mean prediction of each batch $\mathcal{B}$:

$$\mathcal{L}_b(\theta) = \sum_{k=1}^K p_k \log \frac{p_k}{\bar{f}_k(\mathbf{x}, \theta)} = - \sum_{k=1}^K p_k \log \bar{f}_k(\mathbf{x}, \theta),$$

where $\bar{f}_k(\mathbf{x}; \theta) \approx \frac{1}{|\mathcal{B}|} \sum_{\mathbf{x} \in \mathcal{B}} f(\mathbf{x}; \theta)$, and p_k stands for the prior probability of the k th class.

The final loss function for SOP+ is therefore constructed by three terms as follows

$$L(\theta, \{u_i, v_i\}) + \lambda_C \mathcal{L}_C(\theta) + \lambda_B \mathcal{L}_B(\theta),$$

where $\lambda_C, \lambda_B > 0$ are the hyper-parameters.

A.4. Experimental Settings

Data processing: For experiments on CIFAR10/100 (Krizhevsky et al., 2009) without extra techniques, we use simple data augmentations including random crop and horizontal flip following previous works (Patrini et al., 2017; Liu et al., 2020). For SOP+, we use the default setting from unsupervised data augmentation (Xie et al., 2020) to apply efficient data augmentation to create another view of the data for consistency training. For Clothing-1M (Xiao et al., 2015), we first resize images to 256×256 , and then random crop to 224×224 , following a random horizontal flip. For WebVision (Li et al., 2017), we randomly crop the images into size of 227×227 . All images are standardized by their means and variances.

Hyper-parameters of SOP: We adopt a SGD optimizer without weight decay for U and V . We keep all the hyper-parameters fixed for different levels of noise. For fair comparison, we adopt two settings of hyper-parameters and architectures for SOP and SOP+. More details of hyper-parameters can be found in Table A.4. Note that the method is not very sensitive to hyper-parameters λ_C and λ_B .

B. Proofs for Theoretical Analysis with Linear Models

B.1. Proof of Proposition 3.2

We first present a simple but useful lemma.

Lemma B.1. Let $(\theta, \mathbf{u}, \mathbf{v})$ be a critical point to (15) that is not a global minimum, i.e.,

$$\mathbf{r} := \mathbf{J}\theta + \mathbf{u} \odot \mathbf{u} - \mathbf{v} \odot \mathbf{v} - \mathbf{y} \neq \mathbf{0}.$$

Then there exists an index i such that

$$u^i = v^i = 0, \quad r^i \neq 0, \quad (\text{B.1})$$

where u^i , v^i , and r^i denote the i -th elements of $\mathbf{u}$, $\mathbf{v}$ and $\mathbf{r}$, respectively.

Proof. We may compute the gradient of the objective function h in (15) as

$$\begin{aligned} \nabla_{\boldsymbol{\theta}} h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) &= \mathbf{J}^\top \mathbf{r}, \\ \nabla_{\mathbf{u}} h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) &= 2\mathbf{r} \odot \mathbf{u}, \\ \nabla_{\mathbf{v}} h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) &= -2\mathbf{r} \odot \mathbf{v}. \end{aligned}$$

Since $\mathbf{r} \neq \mathbf{0}$ but $\nabla_{\mathbf{u}} h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) = \nabla_{\mathbf{v}} h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) = \mathbf{0}$, we must have $u^i = v^i = 0$ and $r^i \neq 0$ for some i . $\square$

We now prove Proposition 3.2 as follows.

Proof of Proposition 3.2. We compute the hessian $\nabla^2 h$ of the objective function h in (15) as

$$\nabla^2 h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) = \begin{bmatrix} \mathbf{J}^\top \mathbf{J} & 2\mathbf{J}^\top \text{diag}(\mathbf{u}) & -2\mathbf{J}^\top \text{diag}(\mathbf{v}) \\ 2\text{diag}(\mathbf{u})\mathbf{J} & 2\text{diag}(\mathbf{r} + 2\mathbf{u} \odot \mathbf{u}) & -4\text{diag}(\mathbf{v} \odot \mathbf{u}) \\ -2\text{diag}(\mathbf{v})\mathbf{J} & -4\text{diag}(\mathbf{v} \odot \mathbf{u}) & -2\text{diag}(\mathbf{r} - 2\mathbf{v} \odot \mathbf{v}) \end{bmatrix}.$$

For any direction $\mathbf{d} = [\mathbf{d}_\theta^\top \quad \mathbf{d}_u^\top \quad \mathbf{d}_v^\top]^\top$, the quadratic form of the Hessian $\nabla^2 h$ along this direction is given by

$$\begin{aligned} \mathbf{d}^\top \nabla^2 h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) \mathbf{d} &= \|\mathbf{J} \mathbf{d}_\theta\|_2^2 + 4\|\mathbf{u} \odot \mathbf{d}_u\|_2^2 + 4\|\mathbf{v} \odot \mathbf{d}_v\|_2^2 \\ &\quad + 2\langle \mathbf{r}, \mathbf{d}_u \odot \mathbf{d}_u - \mathbf{d}_v \odot \mathbf{d}_v \rangle + 4\langle \mathbf{J} \mathbf{d}_\theta, \mathbf{u} \odot \mathbf{d}_u - \mathbf{v} \odot \mathbf{d}_v \rangle - 8\langle \mathbf{u} \odot \mathbf{d}_u, \mathbf{v} \odot \mathbf{d}_v \rangle. \end{aligned} \quad (\text{B.2})$$

We now consider an arbitrary critical point $(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v})$ of (15) that is not a global minimum. By Lemma B.1, there exists an i such that $r^i \neq 0$ while $u^i = v^i = 0$. We divide the discussion into two cases.

- **Case 1:** $r^i > 0$. We set $\mathbf{d}_\theta = \mathbf{0}$, $\mathbf{d}_u = \mathbf{0}$, and $\mathbf{d}_v$ to be such that all of its entries are zero except for the i -th entry which is given by $d_v^i = 1$. Plugging this direction into (B.2), we obtain

$$\mathbf{d}^\top \nabla^2 h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) \mathbf{d} = 4 \underbrace{[v^i]^2}_{v^i=0} [d_v^i]^2 - 2r^i \underbrace{[d_v^i]^2}_{d_v^i=1} = -2r^i < 0$$

- **Case 2:** $r^i < 0$. We set $\mathbf{d}_\theta = \mathbf{0}$, $\mathbf{d}_v = \mathbf{0}$, and $\mathbf{d}_u$ to be such that all of its entries are zero except for the i -th entry which is given by $d_u^i = 1$. Plugging this direction into (B.2), we obtain

$$\mathbf{d}^\top \nabla^2 h(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v}) \mathbf{d} = 4 \underbrace{[u^i]^2}_{u^i=0} [d_u^i]^2 + 2r^i \underbrace{[d_u^i]^2}_{d_u^i=1} = 2r^i < 0$$

In both cases above we have constructed a direction of negative curvature, hence $(\boldsymbol{\theta}, \mathbf{u}, \mathbf{v})$ is a strict saddle. $\square$

B.2. Proof of Proposition 3.3

The proof is based on the following lemma which follows trivially from KKT conditions:

Lemma B.2 (KKT condition). *Given any $\mathbf{J}$ and $\mathbf{y}$, if there exists $(\widehat{\boldsymbol{\theta}}, \widehat{\mathbf{s}}, \widehat{\mathbf{v}})$ satisfying*

$$\begin{aligned} \mathbf{y} &= \mathbf{J} \widehat{\boldsymbol{\theta}} + \widehat{\mathbf{s}}, \\ \widehat{\boldsymbol{\theta}} &= \mathbf{J}^\top \widehat{\mathbf{v}}, \text{ and} \\ \widehat{\mathbf{v}} &\in \lambda \text{sign}(\widehat{\mathbf{s}}), \end{aligned} \quad (\text{B.3})$$

then $(\widehat{\boldsymbol{\theta}}, \widehat{\mathbf{s}})$ is an optimal solution to (26). In above, $\text{sign}(z)$ is defined entrywise on z as

$$\text{sign}(z) = \begin{cases} z/|z| & \text{if } \widehat{z} \neq 0, \\ [-1, 1] & \text{if } \widehat{z} = 0. \end{cases} \quad (\text{B.4})$$

Proof of Proposition 3.3. We divide the proof into two parts.

Global convergence. In this part we show that $(\boldsymbol{\theta}_\infty(\gamma, \alpha), \mathbf{u}_\infty(\gamma, \alpha), \mathbf{v}_\infty(\gamma, \alpha))$ is a global solution to (15) for any fixed (γ, α) . Denote

$$\mathbf{r}_\infty(\gamma, \alpha) \doteq \lim_{t \rightarrow \infty} \mathbf{r}_t(\gamma, \alpha). \quad (\text{B.5})$$

It follows from (18) and (22) that the limit $\mathbf{r}_\infty(\gamma, \alpha)$ exists and can be written as

$$\mathbf{r}_\infty(\gamma, \alpha) = \mathbf{J}\boldsymbol{\theta}_\infty(\gamma, \alpha) + \mathbf{u}_\infty(\gamma, \alpha) \odot \mathbf{u}_\infty(\gamma, \alpha) - \mathbf{v}_\infty(\gamma, \alpha) \odot \mathbf{v}_\infty(\gamma, \alpha). \quad (\text{B.6})$$

Suppose for the purpose of obtaining a contradiction that $(\boldsymbol{\theta}_\infty(\gamma, \alpha), \mathbf{u}_\infty(\gamma, \alpha), \mathbf{v}_\infty(\gamma, \alpha))$ is not a global solution to (15). It follows from Lemma B.1 that there exists an i such that

$$u_\infty^i(\gamma, \alpha) = v_\infty^i(\gamma, \alpha) = 0, \text{ and } r_\infty^i(\gamma, \alpha) \neq 0. \quad (\text{B.7})$$

Without loss of generality we assume that $C \doteq r_\infty^i(\gamma, \alpha) > 0$ so that $r_t^i(\gamma, \alpha) \rightarrow C$ with $t \rightarrow \infty$. For any $\epsilon \in (0, C)$, there exists a $t_0 > 0$ such that

$$C - \epsilon \leq r_t^i(\gamma, \alpha) \leq C + \epsilon, \quad \forall t > t_0. \quad (\text{B.8})$$

It follows from (B.8) and (21) that

$$\begin{aligned} \nu_t^i(\gamma, \alpha) &= - \int_0^t r_\tau^i(\gamma, \alpha) d\tau = \nu_{t_0}^i(\gamma, \alpha) - \int_{t_0}^t r_\tau^i(\gamma, \alpha) d\tau \\ &\in \left(\nu_{t_0}^i(\gamma, \alpha) - (C + \epsilon)(t - t_0), \nu_{t_0}^i(\gamma, \alpha) - (C - \epsilon)(t - t_0) \right), \quad \forall t > t_0. \end{aligned} \quad (\text{B.9})$$

Using this bound on $\nu_t^i(\gamma, \alpha)$ in (20), we obtain

$$v_t^i(\gamma, \alpha) = \gamma \exp \left(-2\alpha \nu_t^i(\gamma, \alpha) \right) \geq \gamma \exp \left(-2\alpha \nu_{t_0}^i(\gamma, \alpha) \right) \exp \left(2\alpha(C - \epsilon)(t - t_0) \right), \quad \forall t > t_0. \quad (\text{B.10})$$

Taking the limit of $t \rightarrow \infty$, we obtain $v_\infty^i(\gamma, \alpha) = \infty$ which contradicts $v_\infty^i(\gamma, \alpha) = 0$ in (B.7). Therefore, we conclude that $(\boldsymbol{\theta}_\infty(\gamma, \alpha), \mathbf{u}_\infty(\gamma, \alpha), \mathbf{v}_\infty(\gamma, \alpha))$ is a global solution to (15).

Implicit regularization. In this part we prove that $(\hat{\boldsymbol{\theta}}, \hat{\mathbf{s}})$ is an optimal solution to the regularized convex optimization problem in (26). Let $\boldsymbol{\nu}_\infty(\gamma, \alpha)$ be the limit of $\boldsymbol{\nu}_t(\gamma, \alpha)$ in (21) at $t \rightarrow \infty$, and let

$$\hat{\boldsymbol{\nu}} \doteq \lim_{\gamma \rightarrow 0} \boldsymbol{\nu}_\infty(\gamma, \alpha(\gamma)), \quad (\text{B.11})$$

with $\alpha(\gamma)$ defined in (23). We only need to show that the triplet $(\hat{\boldsymbol{\theta}}, \hat{\mathbf{s}}, \hat{\boldsymbol{\nu}})$ with $\hat{\boldsymbol{\theta}}$ defined in (24), $\hat{\mathbf{s}}$ defined in (25) and $\hat{\boldsymbol{\nu}}$ defined in (B.11) satisfies the KKT conditions in (B.3).

1. Because $(\boldsymbol{\theta}_\infty(\gamma, \alpha), \mathbf{u}_\infty(\gamma, \alpha), \mathbf{v}_\infty(\gamma, \alpha))$ is a global solution to (15), we have

$$\mathbf{J}\boldsymbol{\theta}_\infty(\gamma, \alpha) + \mathbf{u}_\infty(\gamma, \alpha) \odot \mathbf{u}_\infty(\gamma, \alpha) - \mathbf{v}_\infty(\gamma, \alpha) \odot \mathbf{v}_\infty(\gamma, \alpha) = \mathbf{y}, \quad \forall \gamma > 0, \alpha > 0. \quad (\text{B.12})$$

Taking the limit of $\gamma \rightarrow 0$ with $\alpha = \alpha(\gamma)$ and noting the assumption that all limits in (24) exist, we obtain

$$\mathbf{J}\hat{\boldsymbol{\theta}} + \hat{\mathbf{u}} \odot \hat{\mathbf{u}} - \hat{\mathbf{v}} \odot \hat{\mathbf{v}} = \mathbf{y}. \quad (\text{B.13})$$

Plugging in the definition of $\hat{\mathbf{s}}$ in (25), we obtain

$$\mathbf{y} = \mathbf{J}\hat{\boldsymbol{\theta}} + \hat{\mathbf{s}}.$$

2. By taking the limit of the relation $\boldsymbol{\theta}_t(\gamma, \alpha) = \mathbf{J}^\top \boldsymbol{\nu}_t(\gamma, \alpha)$ in (20) and noting the assumptions that all relevant limits exist, we obtain

$$\hat{\boldsymbol{\theta}} = \lim_{\gamma \rightarrow 0} \lim_{t \rightarrow \infty} \boldsymbol{\theta}_t(\gamma, \alpha(\gamma)) = \lim_{\gamma \rightarrow 0} \lim_{t \rightarrow \infty} \mathbf{J}^\top \boldsymbol{\nu}_t(\gamma, \alpha(\gamma)) = \mathbf{J}^\top \hat{\boldsymbol{\nu}}. \quad (\text{B.14})$$

3. Denote $\mathbf{s}_\infty(\gamma, \alpha) \doteq \mathbf{u}_\infty(\gamma, \alpha) \odot \mathbf{u}_\infty(\gamma, \alpha) - \mathbf{v}_\infty(\gamma, \alpha) \odot \mathbf{v}_\infty(\gamma, \alpha)$. By (20), we have

$$s_\infty^i(\gamma, \alpha) \doteq u_\infty^i(\gamma, \alpha)^2 - v_\infty^i(\gamma, \alpha)^2 = \gamma^2 \exp(4\alpha \nu_\infty^i(\gamma, \alpha)) - \gamma^2 \exp(-4\alpha \nu_\infty^i(\gamma, \alpha)). \quad (\text{B.15})$$

For each entry of $\hat{\mathbf{s}} = \lim_{\gamma \rightarrow 0} \mathbf{s}_\infty(\gamma, \alpha(\gamma))$ (recall that $\alpha(\gamma)$ is defined in (23)), we may have three cases:

- **Case 1:** $\hat{s}^i > 0$. From (B.15), we must have $\alpha(\gamma) \nu_\infty^i(\gamma, \alpha(\gamma)) \rightarrow +\infty$ as $\gamma \rightarrow 0$ so that

$$\lim_{\gamma \rightarrow 0} \exp \left(4\alpha(\gamma) \nu_\infty^i(\gamma, \alpha(\gamma)) \right) = \infty, \text{ and } \lim_{\gamma \rightarrow 0} \exp \left(-4\alpha(\gamma) \nu_\infty^i(\gamma, \alpha(\gamma)) \right) = 0. \quad (\text{B.16})$$

Hence,

$$\begin{aligned}
 & \lim_{\gamma \rightarrow 0} \gamma^2 \exp \left(4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = \widehat{s}^i \\
 \implies & \lim_{\gamma \rightarrow 0} 2 \log \gamma + 4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) = \log \widehat{s}^i \\
 \implies & \lim_{\gamma \rightarrow 0} \nu_{\infty}^i(\gamma, \alpha(\gamma)) = \lim_{\gamma \rightarrow 0} \left(\frac{\log \widehat{s}^i}{4\alpha(\gamma)} - \frac{\log \gamma}{2\alpha(\gamma)} \right).
 \end{aligned} \tag{B.17}$$

Plugging in the relation $\alpha(\gamma) = -\frac{\log \gamma}{2\lambda}$ in (23), we have

$$\lim_{\gamma \rightarrow 0} \nu_{\infty}^i(\gamma, \alpha(\gamma)) = -\lim_{\gamma \rightarrow 0} \frac{\lambda \log \widehat{s}^i}{2 \log \gamma} + \lambda = \lambda.$$

- **Case 2:** $\widehat{s}^i < 0$. Similar to case 1, from (B.15) we must have

$$\lim_{\gamma \rightarrow 0} \exp \left(4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = 0, \text{ and } \lim_{\gamma \rightarrow 0} \exp \left(-4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = \infty. \tag{B.18}$$

Hence,

$$\begin{aligned}
 & \lim_{\gamma \rightarrow 0} -\gamma^2 \exp \left(-4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = \widehat{s}^i \\
 \implies & \lim_{\gamma \rightarrow 0} 2 \log \gamma - 4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) = \log(-\widehat{s}^i) \\
 \implies & \lim_{\gamma \rightarrow 0} \nu_{\infty}^i(\gamma, \alpha(\gamma)) = \lim_{\gamma \rightarrow 0} \left(-\frac{\log(-\widehat{s}^i)}{4\alpha(\gamma)} + \frac{\log \gamma}{2\alpha(\gamma)} \right).
 \end{aligned} \tag{B.19}$$

Plugging in the relation $\alpha(\gamma) = -\frac{\log \gamma}{2\lambda}$ in (23), we have

$$\lim_{\gamma \rightarrow 0} \nu_{\infty}^i(\gamma, \alpha(\gamma)) = -\lambda.$$

- **Case 3:** $\widehat{s}^i = 0$. From (B.15), we must have

$$\lim_{\gamma \rightarrow 0} \gamma^2 \exp \left(4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = 0, \text{ and } \lim_{\gamma \rightarrow 0} \gamma^2 \exp \left(-4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) = 0. \tag{B.20}$$

Hence, for any small $\epsilon \in (0, 1)$, there exists $\gamma_0 > 0$ such that for all $\gamma \in (0, \gamma_0)$, we have

$$\begin{aligned}
 & \gamma^2 \cdot \max \left\{ \exp \left(4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right), \exp \left(-4\alpha(\gamma) \nu_{\infty}^i(\gamma, \alpha(\gamma)) \right) \right\} < \epsilon \\
 \implies & 2 \log \gamma + 4\alpha(\gamma) \cdot |\nu_{\infty}^i(\gamma, \alpha(\gamma))| < \log \epsilon \\
 \implies & |\nu_{\infty}^i(\gamma, \alpha(\gamma))| < \frac{\log \epsilon}{4\alpha(\gamma)} - \frac{\log \gamma}{2\alpha(\gamma)}.
 \end{aligned} \tag{B.21}$$

Now, plugging $\alpha(\gamma) = -\frac{\log \gamma}{2\lambda}$ in, we have

$$|\nu_{\infty}^i(\gamma, \alpha(\gamma))| < -\frac{\lambda \log \epsilon}{2 \log \gamma} + \lambda < \lambda.$$

Therefore, we have

$$\lim_{\gamma \rightarrow 0} |\nu_{\infty}^i(\gamma, \alpha(\gamma))| < \lambda.$$

Synthesizing all the above three cases, we obtain:

$$\widehat{\mathbf{v}} \in \lambda \text{sign}(\widehat{\mathbf{s}}).$$

□

B.3. Proof of Proposition 3.5

We begin with introducing the null space property that is widely used for providing necessary and sufficient conditions for exact recovery of sparse signals in compressive sensing.

Definition B.3 ((Cohen et al., 2009)). We say a matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$ satisfies the null space property with constant $\rho \in (0, 1)$ relative to $S \in [n]$ if

$$\|\mathbf{v}_S\|_1 \leq \rho \|\mathbf{v}_{S^c}\|_1 \text{ for all } \mathbf{v} \in \ker \mathbf{A},$$

where $\ker \mathbf{A}$ is the null space of $\mathbf{A}$.

With Definition B.3, we prove Proposition 3.5 by using the following two lemmas. The first lemma establishes correct recovery of (θ_*, s_*) from (26) under the null space property.

Lemma B.4. *Given matrix $\mathbf{J}$ and a matrix $\mathbf{A}$ that annihilates $\mathbf{J}$ on the left (i.e. such that $\mathbf{A}\mathbf{J} = 0$). If $\mathbf{A}$ satisfies the stable null space property with constant $\rho \in (0, 1)$ relative to the support of s_* , then the solution to (26) is (θ_*, s_*) for any $\lambda > \lambda_0$ where λ_0 is a scalar that depends only on $(\mathbf{J}, \theta_*, \rho)$.*

The second lemma shows that the null space property is satisfied under the incoherent condition in (28).

Lemma B.5. *Given matrix $\mathbf{J}$ and a matrix $\mathbf{A}$ that annihilates $\mathbf{J}$ on the left, if*

$$k^2 r \leq \frac{N}{\mu(\mathbf{J})} \left(\frac{\rho}{\rho + 1} \right)^2, \quad (\text{B.22})$$

then $\mathbf{A}$ satisfies null space property with constant ρ relative to any S that satisfies $|S| = k$.

Proof of Proposition 3.5. Assume that the condition in (28) is satisfied. Then there exists a $\rho \in (0, 1)$ such that the condition in (B.22) holds. Hence, $\mathbf{A}$ satisfies null space property with constant ρ relative to any S that satisfies $|S| = k$. Since s_* is k -sparse, we have that $\mathbf{A}$ satisfies null space property with constant ρ relative to the support of s_* . Then the conclusion of Proposition 3.5 follows from applying Lemma B.4. Finally, from Lemma B.4 we have that λ_0 is a function of $(\mathbf{J}, \theta_*, \rho)$, wherein ρ is determined by $\mathbf{A}$ (hence $\mathbf{J}$) and the associated sparsity k . Hence λ_0 can be determined with a given $(\mathbf{J}, \theta_*, k)$. $\square$

In the rest of this section we prove Lemma B.4 and Lemma B.5.

Proof of Lemma B.4. We first introduce the following result on a useful property of the stable null space property.

Theorem B.6 (Useful property of stable null space property). *Suppose a matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$ satisfies the null space property with constant $\rho \in (0, 1)$ relative to $S \in [n]$. Then for every vector $\mathbf{x}$ supported on S , we have*

$$\|\mathbf{z} - \mathbf{x}\|_1 \leq \frac{1 + \rho}{1 - \rho} (\|\mathbf{z}\|_1 - \|\mathbf{x}\|_1)$$

for any $\mathbf{z}$ with $\mathbf{A}\mathbf{z} = \mathbf{A}\mathbf{x}$.

Proof of Theorem B.6. Since $\mathbf{A}(\mathbf{z} - \mathbf{x}) = \mathbf{0}$, i.e., $\mathbf{z} - \mathbf{x} \in \ker \mathbf{A}$, the null space property of $\mathbf{A}$ implies

$$\|(\mathbf{z} - \mathbf{x})_{S^c}\|_1 \leq \rho \|(\mathbf{z} - \mathbf{x})_S\|_1,$$

which further gives that

$$\|\mathbf{z} - \mathbf{x}\|_1 \leq (1 + \rho) \|(\mathbf{z} - \mathbf{x})_{S^c}\|_1.$$

We now use these properties to prove the main result as

$$\begin{aligned} \|\mathbf{z}\|_1 &= \|\mathbf{z}_S\|_1 + \|\mathbf{z}_{S^c}\|_1 = \|(\mathbf{z} - \mathbf{x} + \mathbf{x})_S\|_1 + \|\mathbf{z}_{S^c}\|_1 \\ &\geq \|\mathbf{x}\|_1 - \|(\mathbf{z} - \mathbf{x})_S\|_1 + \|(\mathbf{z} - \mathbf{x})_{S^c}\|_1 \\ &\geq \|\mathbf{x}\|_1 + (1 - \rho) \|(\mathbf{z} - \mathbf{x})_{S^c}\|_1 \\ &\geq \|\mathbf{x}\|_1 + \frac{1 + \rho}{1 - \rho} \|\mathbf{z} - \mathbf{x}\|_1, \end{aligned}$$

where the first inequality follows because $\mathbf{x}$ is only supported on S . $\square$

We are now ready to prove Lemma B.4.

Proof of Lemma B.4. Let $\mathbf{J} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$ be the compact SVD of $\mathbf{J}$ and $\mathbf{V}_\perp$ be an orthonormal basis that complements $\mathbf{V}$. Then, the constraint in (26) is equivalent to

$$\mathbf{A}\mathbf{y} = \mathbf{A}\mathbf{s}, \quad \boldsymbol{\theta} = \mathbf{V}\mathbf{\Sigma}^{-1}\mathbf{U}^\top(\mathbf{y} - \mathbf{s}) + \mathbf{V}_\perp\mathbf{h}.$$

Thus, the problem (26) is equivalent to

$$\begin{aligned} \min_{\mathbf{s}, \mathbf{h}} \quad & \frac{1}{2} \|\mathbf{V}\mathbf{\Sigma}^{-1}\mathbf{U}^\top(\mathbf{y} - \mathbf{s}) + \mathbf{V}_\perp\mathbf{h}\|_2^2 + \lambda \|\mathbf{s}\|_1 \\ \text{s.t.} \quad & \mathbf{A}\mathbf{y} = \mathbf{A}\mathbf{s}, \end{aligned} \quad (\text{B.23})$$

which is further equivalent to

$$\begin{aligned} \min_{\mathbf{s}} g(\mathbf{s}) &:= \frac{1}{2} \|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{y} - \mathbf{s})\|_2^2 + \lambda \|\mathbf{s}\|_1 \\ \text{s.t. } \mathbf{A}\mathbf{s}_\star &= \mathbf{A}\mathbf{s}. \end{aligned} \quad (\text{B.24})$$

Assume $\mathbf{A}$ satisfies the stable null space property with constant $\rho \in (0, 1)$ relative to the support of $\mathbf{s}_\star$. Now for any $\mathbf{s}$ with $\mathbf{A}\mathbf{s}_\star = \mathbf{A}\mathbf{s}$, by Theorem B.6, we have

$$\|\mathbf{s}\|_1 - \|\mathbf{s}_\star\|_1 \geq \frac{1-\rho}{1+\rho} \|\mathbf{s} - \mathbf{s}_\star\|_1,$$

which ensures $\mathbf{s} = \mathbf{s}_\star$ if we only minimize $\|\mathbf{s}\|_1$. The first term in (B.24) can be written as

$$\|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{y} - \mathbf{s})\|_2^2 = \|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{s}_\star - \mathbf{s} + \mathbf{J}\boldsymbol{\theta}_\star)\|_2^2,$$

where

$$\boldsymbol{\theta}_\star = \mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{y} - \mathbf{s}_\star).$$

This together with the previous equation gives

$$\begin{aligned} g(\mathbf{s}) - g(\mathbf{s}_\star) &\geq \lambda \frac{1-\rho}{1+\rho} \|\mathbf{s} - \mathbf{s}_\star\|_1 + \|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{s}_\star - \mathbf{s} + \mathbf{J}\boldsymbol{\theta}_\star)\|_2^2 - \|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top\mathbf{J}\boldsymbol{\theta}_\star\|_2^2 \\ &= \lambda \frac{1-\rho}{1+\rho} \|\mathbf{s} - \mathbf{s}_\star\|_1 + \|\mathbf{V}\Sigma^{-1}\mathbf{U}^\top(\mathbf{s}_\star - \mathbf{s})\|_2^2 + 2\langle \mathbf{s}_\star - \mathbf{s}, \mathbf{U}\Sigma^{-1}\mathbf{V}^\top\boldsymbol{\theta}_\star \rangle \\ &\geq \lambda \frac{1-\rho}{1+\rho} \|\mathbf{s} - \mathbf{s}_\star\|_1 - 2\|\mathbf{U}\Sigma^{-1}\mathbf{V}^\top\boldsymbol{\theta}_\star\|_\infty \|\mathbf{s} - \mathbf{s}_\star\|_1 \\ &= \left(\lambda \frac{1-\rho}{1+\rho} - 2\|\mathbf{U}\Sigma^{-1}\mathbf{V}^\top\boldsymbol{\theta}_\star\|_\infty \right) \|\mathbf{s} - \mathbf{s}_\star\|_1. \end{aligned}$$

Thus, if $\lambda > \lambda_0$ with

$$\lambda_0 = 2 \frac{1+\rho}{1-\rho} \|\mathbf{U}\Sigma^{-1}\mathbf{V}^\top\boldsymbol{\theta}_\star\|_\infty, \quad (\text{B.25})$$

we have $g(\mathbf{s}) - g(\mathbf{s}_\star) > 0$ whenever $\mathbf{s} \neq \mathbf{s}_\star$. □

Proof of Lemma B.5.

Proof. Let $\mathbf{J} = \mathbf{U}\Sigma\mathbf{V}^\top$ be the compact SVD of $\mathbf{J}$. From (B.22) we have

$$(\rho + 1)k\sqrt{\frac{r}{N}\mu(\mathbf{J})} \leq \rho. \quad (\text{B.26})$$

Let $S \subseteq [N]$ with $|S| = k$ and $\mathbf{a} \in \mathbb{R}^r$ be an arbitrary vector. We have

$$\|[\mathbf{U}\mathbf{a}]_S\|_1 = \sum_{i \in S} |\mathbf{e}_i^\top \mathbf{U}\mathbf{a}| = \sum_{i \in S} |\langle \mathbf{U}^\top \mathbf{e}_i, \mathbf{a} \rangle| \leq k\|\mathbf{a}\|_2 \cdot \max_{i \in S} \|\mathbf{U}^\top \mathbf{e}_i\|_2 \leq k\|\mathbf{a}\|_2 \sqrt{\frac{r}{N}\mu(\mathbf{J})}, \quad (\text{B.27})$$

where the last inequality is obtained from Definition 3.4. In addition, we have

$$\|\mathbf{U}\mathbf{a}\|_1 \geq \|\mathbf{U}\mathbf{a}\|_2 = \|\mathbf{a}\|_2. \quad (\text{B.28})$$

Combining (B.26), (B.27) and (B.28), we get

$$(\rho + 1)\|[\mathbf{U}\mathbf{a}]_S\|_1 \leq (\rho + 1)k\|\mathbf{a}\|_2 \sqrt{\frac{r}{N}\mu(\mathbf{J})} \leq \rho\|\mathbf{a}\|_2 \leq \rho\|\mathbf{U}\mathbf{a}\|_1, \quad (\text{B.29})$$

hence,

$$\|[\mathbf{U}\mathbf{a}]_S\|_1 \leq \rho\|[\mathbf{U}\mathbf{a}]_{S^c}\|_1. \quad (\text{B.30})$$

Noting that $\{\mathbf{U}\mathbf{a} | \mathbf{a} \in \mathbb{R}^r\} = \ker \mathbf{A}$, this finishes the proof by Definition B.3. □