

# Training (Overparametrized) Neural Networks in Near-Linear Time

## Abstract

The slow convergence rate and pathological curvature issues of first-order gradient methods for training deep neural networks, initiated an ongoing effort for developing faster *second-order* optimization algorithms beyond SGD, without compromising the generalization error. Despite their remarkable convergence rate (*independent* of the training batch size  $n$ ), second-order algorithms incur a daunting slowdown in the *cost per iteration* (inverting the Hessian matrix of the loss function), which renders them impractical. Very recently, this computational overhead was mitigated by the works of [ZMG19, CGH<sup>+</sup>19], yielding an  $O(mn^2)$ -time second-order algorithm for training two-layer overparametrized neural networks of polynomial width  $m$ .

We show how to speed up the algorithm of [CGH<sup>+</sup>19], achieving an  $\tilde{O}(mn)$ -time backpropagation algorithm for training (mildly overparametrized) ReLU networks, which is near-linear in the dimension  $(mn)$  of the full gradient (Jacobian) matrix. The centerpiece of our algorithm is to reformulate the Gauss-Newton iteration as an  $\ell_2$ -regression problem, and then use a Fast-JL type dimension reduction to *precondition* the underlying Gram matrix in time independent of  $M$ , allowing to find a sufficiently good approximate solution via *first-order* conjugate gradient. Our result provides a proof-of-concept that advanced machinery from randomized linear algebra—which led to recent breakthroughs in *convex optimization* (ERM, LPs, Regression)—can be carried over to the realm of deep learning as well.

# 1 Introduction

Understanding the dynamics of gradient-based optimization of deep neural networks has been a central focal point of theoretical machine learning in recent years [LY17, ZSJ<sup>+</sup>17, ZSD17, LL18, DZPS19, AZLS19a, AZLS19b, AZLL19, BJW19, OS19, ADH<sup>+</sup>19b, CG19, SY19, Dan20, JT20, BELM20]. This line of work led to a remarkable rigorous understanding of the generalization, robustness and convergence rate of *first-order* (SGD-based) algorithms, which are the standard choice for training DNNs. By contrast, the *computational complexity* of implementing gradient-based training algorithms (e.g., backpropagation) in such non-convex landscape is less understood, and gained traction only recently due to the overwhelming size of training data and complexity of network design [MG15, DHS11, LJH<sup>+</sup>19, CGH<sup>+</sup>19, ZMG19].

The widespread use first-order methods such as (stochastic) gradient descent in training DNNs is explained, to a large extent, by its computational efficiency – recalculating the gradient of the loss function at each iteration is simple and cheap (linear in the dimension of the full gradient), let alone with the advent of minibatch random sampling [HRS16, CGH<sup>+</sup>19]. Nevertheless, first-order methods have a slow rate of convergence in non-convex settings (typically  $\Omega(\text{poly}(n) \log(1/\epsilon))$  for overparametrized networks, see e.g., [ZMG19]) for reducing the training error below  $\epsilon$ , and it is increasingly clear that SGD-based algorithms are becoming a real bottleneck for many practical purposes. This drawback initiated a substantial effort for developing fast training methods beyond SGD, aiming to improve its convergence rate without compromising the generalization error [BLC88, Mar10, MG15, DHS11, KB15, PW17, CGH<sup>+</sup>19, ZMG19].

Second-order gradient algorithms (which employ information about the Hessian of the loss function), pose an intriguing computational tradeoff in this context: On one hand, they are known to converge extremely fast, at a rate *independent* of the input size (i.e., only  $O(\log 1/\epsilon)$  iterations [ZMG19]), and offer a qualitative advantage in overcoming pathological curvature issues that arise in first-order methods, by exploiting the local geometry of the loss function. This feature implies another practical advantage of second order methods, namely, that they do not require tuning the learning rate [CGH<sup>+</sup>19, ZMG19]. On the other hand, second-order methods have a prohibitive *cost per iteration*, as they involve *inverting* a dynamically-changing dense Hessian matrix. This drawback explains the scarcity of second order methods in *large scale non-convex* optimization, in contrast to its popularity in the convex setting.

The recent works of [CGH<sup>+</sup>19, ZMG19] addressed the computational bottleneck of second-order algorithms in optimizing deep neural nets, and presented a training algorithm for overparametrized neural networks with smooth (resp. ReLU) activations, whose running time is  $O(mn^2)$ , where  $m$  is the width of the neural network, and  $n$  is the size of the training data in  $\mathbb{R}^d$ . The two algorithms, which achieve essentially the same running time, are based on the classic Gauss-Newton algorithm (resp. ‘Natural gradient’ algorithm) combined with the recent introduction of *Neural Tangent Kernels* (NTK) [JGH18]. The NTK formulation utilizes a local-linearization of the loss function for overparametrized neural networks, which reduces the optimization problem of DNNs to that of a *kernel regression* problem: The main insight is that when the network is *overparametrized*, i.e., sufficiently wide  $m \gtrsim n^4$  ([SY19]), the neural network becomes locally convex and smooth, hence the problem is equivalent to a kernel regression problem with respect to the NTK function [JGH18], and therefore solving the latter via (S)GD is guaranteed to converge to a global minimum. The training algorithm of [CGH<sup>+</sup>19] draws upon this equivalence, by designing a second-order variation of the Gauss-Newton algorithm (termed ‘Gram-Gauss-Newton’), yielding the aforementioned runtime for *smooth activation functions*.

**Single vs. Multilayer Network Training** Following [CGH<sup>+</sup>19, ZMG19], we focus on two-layer (i.e., single hidden-layer) neural networks. While our algorithm extends to the multilayer case (with a slight compromise on the width dependence), we argue that, as far as training time, the two-layer case is not only the common case, but in fact the *only* interesting case for constant training error: Indeed, in the multilayer case ( $L \geq 2$ ), we claim that the mere cost of *feed-forward* computation of the network’s output is already  $\Omega_\epsilon(m^2 n L)$ . Indeed, the total number of parameters of  $L$ -layer networks is  $M = (L - 1)m^2 + md$ , and as such, feed-forward computation requires, at the very least, computing a single product of  $m \times m$  (dense) matrices  $W$  with a  $m \times 1$  vector for each training data, which already costs  $m^2 n$  time:

$$\hat{y}_i = a^\top \sigma_L \left( \underbrace{W_L}_{m \times m} \sigma_{L-1} \left( \underbrace{W_{L-1}}_{m \times m} \dots \sigma_1 \left( \underbrace{W_1}_{m \times d} x_i \right) \right) \right)$$

Therefore, sublinear-time techniques (as we present) appear futile in the case of multi-layer overparametrized networks, where it is possible to achieve linear time (in  $M$ ) using essentially direct (lossless) computation (see next subsection). It may still be possible to use sublinear algorithms to improve the running time to  $O(m^2 n L + \text{poly}(n))$ , though in for overparametrized DNNs this seems a minor saving.

## 1.1 Our Result

Our main result is a quadratic speedup to the algorithm of [CGH<sup>+</sup>19], yielding an *essentially optimal* training algorithm for overparametrized two-layer neural networks. Moreover, in contrast to [CGH<sup>+</sup>19], our algorithm applies to the more complex and realistic case of *ReLU* activation functions. Our main result is shown below (For a more comprehensive comparison, see Table 1 below and references therein).

**Theorem 1.1.** *Suppose the width of a ReLU neural network satisfies*

$$m = \Omega(\max\{\lambda^{-4} n^4, \lambda^{-2} n^2 d \log(n/\delta)\}),$$

*where  $\lambda > 0$  denotes the minimum eigenvalue of the Gram matrix (see Eq. (5) below). Then with probability  $1 - \delta$  over the random initialization of neural network and the randomness of the training algorithm, our algorithm achieves*

$$\|f_{t+1} - y\|_2 \leq \frac{1}{2} \|f_t - y\|_2.$$

*The computational cost of each iteration is  $\tilde{O}(mnd + n^3)$ , and the running time for reducing the training loss to  $\epsilon$  is  $\tilde{O}((mnd + n^3) \log(1/\epsilon))$ . Using fast matrix-multiplication, the total running time can be further reduced to  $\tilde{O}((mnd + n^\omega) \log(1/\epsilon))$ .<sup>1</sup>*

**Remark 1.2.** *We stress that that our algorithm runs in (near) linear time even for networks with width  $m \gtrsim n^2$  and in fact, under the common belief that  $\omega = 2$ , this is true so long as  $m \gtrsim n$  (!). This means that the bottleneck for linear-time training of small-width DNNs is not computational, but rather analytic: The overparametrization requirements ( $m \gtrsim n^4$ ) in Theorem 1.1 stems from current-best analysis of the convergence guarantees of (S)GD-based training of ReLU networks, and any improvement on these bounds would directly yield linear-time training for thinner networks using our algorithm.*

---

<sup>1</sup>Here,  $\omega < 2.373$  denotes the fast matrix-multiplication (FMM) constant for multiplying two  $n \times n$  matrices [Wil12, LG14].

| Reference             | Method                     | #Iters                     | Cost/iter | Width            | ReLU? |
|-----------------------|----------------------------|----------------------------|-----------|------------------|-------|
| [DZPS19]              | Gradient descent           | $O(n^2 \log(1/\epsilon))$  | $O(mn)$   | $\Omega(n^6)$    | Yes   |
| [SY19]                | Gradient descent           | $O(n^2 \log(1/\epsilon))$  | $O(mn)$   | $\Omega(n^4)$    | Yes   |
| [WDW19]               | Adaptive gradient descent  | $O(n \log(1/\epsilon))$    | $O(mn)$   | $\Omega(n^6)$    | Yes   |
| [CGH <sup>+</sup> 19] | Gram-Gaussian-Newton (GGN) | $O(\log \log(1/\epsilon))$ | $O(mn^2)$ | $\Omega(n^4)$    | No    |
| [CGH <sup>+</sup> 19] | Batch-GGN                  | $O(n^2 \log(1/\epsilon))$  | $O(m)$    | $\Omega(n^{18})$ | No    |
| [ZMG19]               | Natural gradient descent   | $O(\log(1/\epsilon))$      | $O(mn^2)$ | $\Omega(n^4)$    | Yes   |
| <b>This Work</b>      |                            | $O(\log(1/\epsilon))$      | $O(mn)$   | $\Omega(n^4)$    | Yes   |

Table 1: Summary of state-of-art algorithms for training two-layer neural networks.  $n$  denotes the training batch size (number of input data points in  $\mathbb{R}^d$ ) and  $\epsilon$  denote the desired accuracy of the training loss. For simplicity, here we assume  $d = O(1)$  and omit  $\text{poly}(\log n, 1/\lambda)$  terms. The result of [CGH<sup>+</sup>19] applies only to smooth activation gates and not to ReLU networks. Comparison to SGD algorithms is omitted from this table since they require a must stronger assumption on the width  $m$  for convergence, and have slower convergence rate than GD [LL18, AZLS19a, AZLS19b, ZG19].

**Techniques** The majority of ML optimization literature on overparametrized network training is dedicated to understanding and minimizing the *number of iterations* of the training process [ZMG19, CGH<sup>+</sup>19, ZG19] as opposed to the *cost per iteration*, which is the focus of our paper. Our work shows that it is possible to harness the toolbox of *randomized linear algebra*— which was heavily used in the past decade to reduce the cost of *convex optimization* tasks— in the nonconvex setting of deep learning as well. A key ingredient in our algorithm is *linear sketching*, where the main idea is to carefully *compress* a linear system underlying an optimization problem, in a way that preserves a good enough solution to the problem yet can be solved much faster in lower dimension. This is the essence of the celebrated *Sketch-and-Solve* (S&S) paradigm [CW13]. As we explain below, our main *departure* from the classic S&S framework (e.g., [PW17]) is that we cannot afford to directly solve the underlying compressed regression problem (as this approach turns out to be prohibitively slow for our application). Instead, we use sketching (or sampling) to facilitate *fast preconditioning* of linear systems (in the spirit of [ST04, KOSZ13, RT08, Woo14]), which in turn enables to solve the compressed regression problem to very high accuracy via first-order *conjugate* gradient descent. This approach essentially *decouples* the sketching error from the final precision error of the Gauss-Newton step, enabling a much smaller sketch size. We believe this (somewhat unconventional) approach to non-convex optimization is the most enduring message of our work.

## 1.2 Related Work

**Second-order methods in non-convex optimization** Despite the prevalence of first order methods in deep learning applications, there is a vast body of ongoing work [BRB17, BLH18, MG15, GM16, GKS18, CGH<sup>+</sup>19, ZMG19] aiming to design more scalable second-order algorithms that overcome the limitations of (S)GD for optimizing deep models. Grosse and Martens [MG15, GM16] designed the K-FAC method, where the idea is to use Kronecker-factors to approximate the Fisher information matrix, combined with natural gradient descent. This approach has been further explored and extended by [WMG<sup>+</sup>17, GLB<sup>+</sup>18, MBJ18]. Gupta et al. [GKS18] designed the “Shampoo method”, based on the idea of *structure-aware preconditioning*. Anil et al. [AGK<sup>+</sup>20] further validate the practical performance of Shampoo and incorporated it into hardware. However, despite sporadic empirical evidence of such second-order methods (e.g., K-FAC and Shampoo), these methods generally lack a provable theoretical guarantee on the performance when applied to deep

neural networks. Furthermore, in the overparametrized setting, their cost per-iteration in general is at least  $\Omega(mn^2)$ .

We remark that in the *convex* setting, theoretical guarantees for large-scale second-order algorithms have been established (e.g., [ABH17, PW17, MNJ16, ?]), but such rigorous analysis in non-convex setting was only recently proposed ([CGH<sup>+</sup>19, ZMG19]). Our algorithm bears some similarities to the *NewtonSketch* algorithm of [PW17], which also incorporates sketching into second order Newton methods. A key difference, however, is that the algorithm of [PW17] works only for convex problems, and requires access to  $(\nabla^2 f(x))^{1/2}$  (i.e., the square-root of the Hessian). Most importantly, though, [PW17] use the standard (black-box) Sketch-and-Solve paradigm to reduce the computational cost, while this approach incurs large computation overhead in our non-convex setting. By contrast, we use sketching as a subroutine for fast preconditioning. As a by-product, in Section D we show how to apply our techniques to give a substantial improvement over [PW17] in the *convex* setting.

The aforementioned works of [ZMG19] and [CGH<sup>+</sup>19] are most similar in spirit to ours. Zhang et al. [ZMG19] analyzed the convergence rate of Natural gradient descent algorithms for two-layer (overparametrized) neural networks, and showed that the number of iterations is *independent* of the training data size  $n$  (essentially  $\log(1/\epsilon)$ ). They also demonstrate similar results for the convergence rate of K-FAC in the overparametrized regime, albeit with larger requirement on the width  $m$ . Another downside of K-FAC is the high cost per iteration ( $\sim mn^2$ ). Cai et al. [CGH<sup>+</sup>19] analyzed the convergence rate of the so-called Gram-Gauss-Newton algorithm for training two-layer (overparametrized) neural network with *smooth* activation gates. They proved a quadratic (i.e., doubly-logarithmic) convergence rate in this setting ( $\log(\log(1/\epsilon))$ ) albeit with  $O(mn^2)$  cost per iteration. It is noteworthy that this quadratic convergence rate analysis does not readily extend to the more complex and realistic setting of ReLU activation gates, which is the focus of our work. [CGH<sup>+</sup>19] also prove bounds on the convergence of ‘batch GGN’, showing that it is possible to reduce the cost-per-iteration to  $m$ , at the price of  $O(n^2 \log(1/\epsilon))$  iterations, for very heavily overparametrized DNNs (currently  $m = \Omega(n^{18})$ ).

**Sketching** The celebrated ‘Sketch and Solve’ (S&S) paradigm [CW13] was originally developed to speed up the cost of solving linear regression and low-rank approximation problems. This dimensionality-reduction technique has since then been widely developed and applied to both convex and non-convex numerical linear algebra problems [BWZ16, RSW16, WZ16, ALS<sup>+</sup>18, BW18, BCW19, WW19, DJS<sup>+</sup>19, SWY<sup>+</sup>19, Son19, BWZ20], as well as machine-learning applications [AKM<sup>+</sup>17, AKM<sup>+</sup>19, LPPW20, WZ20]. The most direct application of the sketch-and-solve technique is overconstrained regression problems, where the input is a linear system  $[A, b] \in \mathbb{R}^{n \times (d+1)}$  with  $n \gg d$ , and we aim to find an (approximate) solution  $\hat{x} \in \mathbb{R}^d$  so as to minimize the residual error  $\|A\hat{x} - b\|_2$ .

In the classic S&S paradigm, the underlying regression solver is treated as a *black box*, and the computational savings comes from applying it on a smaller *compressed* matrix. Since then, sketching (or sampling) has also been used in a non-black-box fashion for speeding-up optimization tasks, e.g., as a subroutine for preconditioning [Woo14, RT08, ST04, KOSZ13] or fast inverse-maintenance in Linear Programming solvers, semi-definite programming, cutting plane methods, and empirical-risk minimization [CLS19, JSWZ20, JKL<sup>+</sup>20, JLSW20, LSZ19].

**Overparametrization in neural networks** A long and active line of work in recent deep learning literature has focused on obtaining rigorous bounds on the convergence rate of various local-search algorithms for optimizing DNNs [LL18, DZPS19, AZLS19a, AZLS19b, ADH<sup>+</sup>19a, ADH<sup>+</sup>19b,

ZG19, CG19, SY19, JT20]. The breakthrough work of Jacob et al. [JGH18] and subsequent developments<sup>2</sup> introduced the notion of *neural tangent kernels* (NTK), implying that for wide enough networks ( $m \gtrsim n^4$ ), (stochastic) gradient descent provably converges to an optimal solution, with generalization error independent of the number of network parameters.

## 2 Technical Overview

We now provide a streamlined overview of our main result, Theorem 1.1. As discussed in the introduction, our algorithm extends to multi-layer ReLU networks, though we focus on the two-layer case (one-hidden layer), which is the most interesting case where one can indeed hope for linear training time.

The main, and most expensive step, of the GGN (or natural gradient descent) algorithms [CGH<sup>+</sup>19, ZMG19] is multiplying, in each iteration  $t$ , the *inverse* of the Gram matrix  $G_t := J_t J_t^\top$  with the Jacobian matrix  $J_t \in \mathbb{R}^{n \times m}$ , whose  $i$ th row contains the gradient of the  $m = md$  network gates w.r.t the  $i$ th datapoint  $x_i$  (in our case, under ReLU activation).

Naïvely computing  $G_t$  would already take  $mdn^2$  time, however, the *tensor product* structure of the Jacobian  $J$  in fact allows to compute  $G_t$  in  $n \cdot \mathcal{T}_{\text{mat}}(m, d, n) \ll mn^2$  time, where  $\mathcal{T}_{\text{mat}}(m, d, n)$  is the cost of fast rectangular matrix multiplication [Wil12, LG14, GU18].<sup>3</sup> Since the Gram-Gauss-Newton (GGN) algorithm requires  $O(\log \log 1/\epsilon)$  iterations to converge to an  $\epsilon$ -global minimum of the  $\ell_2$  loss [CGH<sup>+</sup>19], this observation yields an  $O(n \cdot \mathcal{T}_{\text{mat}}(m, d, n) \log \log 1/\epsilon)$  total time algorithm for reducing the training loss below  $\epsilon$ . While already nontrivial, this is still far from linear running time ( $\gg mdn$ ).

We show how to carry out each Gauss-Newton iteration in time  $\tilde{O}(mnd + n^3)$ , at the price of slightly compromising the number of iterations to  $O(\log 1/\epsilon)$ , which is inconsequential for the natural regime of constant dimension  $d$  and constant  $\epsilon$ <sup>4</sup>. Our first key step is to reformulate the Gauss-Newton iteration (multiplying  $G_t^{-1}$  by the error vector) as an  $\ell_2$ -regression problem:

$$\min_{g_t} \|J_t J_t^\top g_t - (f_t - y)\|_2 \quad (1)$$

where  $(f_t - y)$  is the training error with respect to the network's output and the training labels  $y$ . Since the Gauss-Newton method is robust to small perturbation errors (essentially [Vai89b, Vai89a]), our analysis shows that it is sufficient to find an approximate solution  $g'_t$  such that  $J_t^\top g'_t$  satisfies

$$\|J_t J_t^\top g'_t - y\|_2 \leq \gamma \|y\|_2, \quad \text{for } \gamma \approx 1/n. \quad (2)$$

The benefit of this reformulation is that it allows to use *linear sketching* to first compress the linear system, significantly reducing the dimension of the optimization problem and thereby the cost of finding a solution, at the price of a small error in the found solution (this is the essence of the *sketch-and-solve* paradigm [CW13]). Indeed, a (variation of) the *Fast-JL* sketch [AC06, LDFU13] guarantees that we can multiply the matrix  $J_t^\top \in \mathbb{R}^{m \times n}$  by a much smaller  $\tilde{O}(n/\delta^2) \times m$  matrix  $S$ , such that (i) the multiplication takes near-linear time  $\tilde{O}(mn)$  time (using the FFT algorithm),

<sup>2</sup>For a complete list of references, we refer the readers to [ADH<sup>+</sup>19a, ADH<sup>+</sup>19b].

<sup>3</sup>To see this, observe that the kronecker-product structure of  $J$  (here  $J \in \mathbb{R}^{n \times md}$  can be constructed from an  $n \times m$  matrix and an  $n \times d$  matrix) allows computing  $Jh$  for any  $h \in \mathbb{R}^{md}$  using fast rectangular matrix multiplication in time  $\mathcal{T}_{\text{mat}}(m, d, n)$  which is near linear time in the dimension of  $J$  and  $h$  (that is,  $n \times m + n \times d$  for  $J$  and  $md$  for  $h$ ) so long as  $d \leq n^\alpha = n^{0.31}$  [GU18], hence computing  $G = JJ^\top$  can be done using  $n$  independent invocations of the aforementioned subroutine, yielding  $n \cdot \mathcal{T}_{\text{mat}}(m, d, n)$  as claimed.

<sup>4</sup>We also remark that this slowdown in the convergence rate is also a consequence of a direct extension of the analysis in [CGH<sup>+</sup>19] to ReLU activation functions.

and (ii)  $SJ_t^\top$  is a  $\delta$ -spectral approximation of  $J_t^\top$  (i.e.,  $\|J_t S^\top S J_t^\top x\|_2 = (1 \pm \delta)\|G_t x\|_2$  for every  $x$ ). Since both computing and inverting the matrix  $\tilde{G}_t := J_t S^\top S J_t^\top$  takes  $\tilde{O}(n^3/\delta^2)$  time, the overall cost of finding a  $\delta$ -approximate solution to the regression problem becomes at most  $\tilde{O}(mn + n^3/\delta^2)$ . Alas, as noted in Equation (2), the approximation error of the found solution must be *polynomially small*  $\gamma \sim 1/n$  in order to guarantee the desired convergence rate (i.e., constant decrease in training error per iteration). This means that we must set  $\delta \sim \gamma \sim 1/n$ , hence the cost of the naïve “sketch-and-solve” algorithm would be at least  $\tilde{O}(n^3/\delta^2) = \tilde{O}(n^5)$ , which is a prohibitively large overhead in both theory and practice (and in particular, no longer yields linear runtime whenever  $m \ll n^4$  which is the current best overparametrization guarantee [SY19]). Since the  $O(1/\delta^2)$  dependence of the JL embedding is known to be tight in general [LN17], this means we need to take a more clever approach to solve the regression (1). This is where our algorithm departs from the naïve sketch-and-solve method, and is the heart of our work.

Our key idea is to use dimension reduction—not to directly invert the compressed matrix—but rather to *precondition* it quickly. More precisely, our approach is to use a (conjugate) gradient-descent solver for the regression problem itself, with a fast preconditioning step, ensuring exponentially faster convergence to very high (polynomially small) accuracy. Indeed, conjugate gradient descent is guaranteed to find a  $\gamma$ -approximate solution to a regression problem  $\min_x \|Ax - b\|_2$  in  $O(\sqrt{\kappa} \log(1/\gamma))$  iterations, where  $\kappa(A)$  is the *condition number* of  $A$  (i.e., the ratio of maximum to minimum eigenvalue). Therefore, if we can ensure that  $\kappa(G_t)$  is small, then we can  $\gamma$ -solve the regression problem in  $\sim mn \log(1/\gamma) = \tilde{O}(mn)$  time, since the per-iteration cost of first-order SGD is linear ( $\sim mn$ ).

The crucial advantage of our approach is that it *decouples* the sketching error from the final precision of the regression problem: Unlike the usual ‘sketch-and-solve’ method, where the sketching error  $\delta$  directly affects the overall precision of the solution to (2), here  $\delta$  only affects the *quality of the preconditioner* (i.e., the ratio of max/min singular values of the sketch  $\tilde{G}_t$ ), hence it suffices to take a *constant* sketching error  $\delta = 0.1$  (say), while letting the SGD deal with the final precision (at it has logarithmic dependence on  $\gamma$ ). See Lemma B.1 for the formal details.

Indeed, by setting the sketching error to  $\delta = 0.1$  (say), the resulting matrix  $\tilde{G}_t = J_t S^\top S J_t^\top$  is small enough ( $n \times \tilde{O}(n)$ ) that we can afford running a standard (QR) algorithm to precondition it, at another  $\tilde{O}(n^3)$  cost per iteration. The output of this step is a matrix  $\tilde{G}'_t := \text{Prec}(\tilde{G}_t)$  with a *constant* condition number  $\kappa(\tilde{G}'_t)$  which preserves  $\tilde{G}'_t x \approx_{\ell_2} \tilde{G}_t x$  up to  $(1 \pm \delta)^2$  relative error. At this point, we can run a (conjugate) gradient descent algorithm, which is guaranteed to find a  $\gamma \approx 1/n$  approximate solution to (1) in time  $\tilde{O}((mn \log((1 + \delta)/\gamma) + n^3))$ , as desired.

We remark that, by definition, the preconditioning step (on the JL sketch) does *not* preserve the eigen-spectrum of  $G_t$ , which is in fact necessary to guarantee the fast convergence of the Gauss-Newton iteration (see Lemma C.3). The point is that this preconditioning step is only preformed as a *local subroutine* so as to solve the regression problem, and does *not* affect the convergence rate of the outer loop.

## 3 Preliminaries

### 3.1 Model and Problem Setup

We denote by  $n$  the number of data points in the training batch, and by  $d$  the data dimension/feature-space (i.e.,  $x_i \in \mathbb{R}^d$ ). We denote by  $m$  the *width* of neural network, and by  $L$  the number of layers and by  $M$  the number of parameters. We assume the data has been normalized, i.e.,  $\|x\|_2 = 1$ . We begin with the two-layer neural network in the following section, and then extend to multilayer networks. Consider a two-layer ReLU activated neural network with  $m$  neurons in the (single) hidden

layer:

$$f(W, x, a) = \frac{1}{\sqrt{m}} \sum_{r=1}^m a_r \phi(w_r^\top x),$$

where  $x \in \mathbb{R}^d$  is the input,  $w_1, \dots, w_m \in \mathbb{R}^d$  are weight vectors in the first layer,  $a_1, \dots, a_m \in \mathbb{R}$  are weights in the second layer. For simplicity, we consider  $a \in \{-1, +1\}^m$  is fixed over all the iterations, this is natural in deep learning theory [LL18, DZPS19, AZLS19a, AZLL19, SY19]. Recall the ReLU function  $\phi(x) = \max\{x, 0\}$ . Therefore for  $r \in [m]$ , we have

$$\frac{\partial f(W, x, a)}{\partial w_r} = \frac{1}{\sqrt{m}} a_r x \mathbf{1}_{w_r^\top x \geq 0}. \quad (3)$$

Given  $n$  input data points  $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) \in \mathbb{R}^d \times \mathbb{R}$ . We define the objective function  $\mathcal{L}$  as follows

$$\mathcal{L}(W) = \frac{1}{2} \sum_{i=1}^n (y_i - f(W, x_i, a))^2.$$

We can compute the gradient of  $\mathcal{L}$  in terms of  $w_r$

$$\frac{\partial \mathcal{L}(W)}{\partial w_r} = \frac{1}{\sqrt{m}} \sum_{i=1}^n (f(W, x_i, a) - y_i) a_r x_i \mathbf{1}_{w_r^\top x_i \geq 0}. \quad (4)$$

We define the prediction function  $f_t : \mathbb{R}^{d \times n} \rightarrow \mathbb{R}^n$  at time  $t$  as follow

$$f_t = \begin{bmatrix} \frac{1}{\sqrt{m}} \sum_{r=1}^m a_r \cdot \phi(\langle w_r(t), x_1 \rangle) \\ \frac{1}{\sqrt{m}} \sum_{r=1}^m a_r \cdot \phi(\langle w_r(t), x_2 \rangle) \\ \vdots \\ \frac{1}{\sqrt{m}} \sum_{r=1}^m a_r \cdot \phi(\langle w_r(t), x_n \rangle) \end{bmatrix}$$

where  $W_t = [w_1(t)^\top, w_2(t)^\top, \dots, w_m(t)^\top]^\top \in \mathbb{R}^{md}$  and  $X = [x_1, x_2, \dots, x_n] \in \mathbb{R}^{d \times n}$ .

For each time  $t$ , the Jacobian matrix  $J \in \mathbb{R}^{n \times md}$  is defined via the following formulation:

$$J_t = \frac{1}{\sqrt{m}} \begin{bmatrix} a_1 x_1^\top \mathbf{1}_{\langle w_1(t), x_1 \rangle \geq 0} & a_2 x_1^\top \mathbf{1}_{\langle w_2(t), x_1 \rangle \geq 0} & \cdots & a_m x_1^\top \mathbf{1}_{\langle w_m(t), x_1 \rangle \geq 0} \\ a_1 x_2^\top \mathbf{1}_{\langle w_1(t), x_2 \rangle \geq 0} & a_2 x_2^\top \mathbf{1}_{\langle w_2(t), x_2 \rangle \geq 0} & \cdots & a_m x_2^\top \mathbf{1}_{\langle w_m(t), x_2 \rangle \geq 0} \\ \vdots & \vdots & \ddots & \vdots \\ a_1 x_n^\top \mathbf{1}_{\langle w_1(t), x_n \rangle \geq 0} & a_2 x_n^\top \mathbf{1}_{\langle w_2(t), x_n \rangle \geq 0} & \cdots & a_m x_n^\top \mathbf{1}_{\langle w_m(t), x_n \rangle \geq 0} \end{bmatrix}.$$

The Gram matrix  $G_t$  is defined as  $G_t = J_t J_t^\top$ , whose  $(i, j)$ -th entry is  $\left\langle \frac{f(W_t, x_i)}{\partial W}, \frac{f(W_t, x_j)}{\partial W} \right\rangle$ . The crucial observation of [JGH18, DZPS19] is that the asymptotic of the Gram matrix equals a positive semidefinite kernel matrix  $K \in \mathbb{R}^{n \times n}$ , where

$$K(x_i, x_j) = \mathbb{E}_{w \in \mathcal{N}(0, 1)} \left[ x_i^\top x_j \mathbf{1}_{\langle w, x_i \rangle \geq 0, \langle w, x_j \rangle \geq 0} \right]. \quad (5)$$

**Assumption 3.1.** We assume the least eigenvalue  $\lambda$  of the kernel matrix  $K$  defined in Eq. (5) satisfies  $\lambda > 0$ .

### 3.2 Subspace embedding

Subspace embedding was first introduced by Sarlós [Sar06], it has been extensively used in numerical linear algebra field over the last decade [CW13, NN13, BW14, SWZ19c]. For a more detailed survey, we refer the readers to [Woo14]. The formal definition is:

**Definition 3.2** (Approximate subspace embedding, ASE [Sar06]). *A  $(1 \pm \epsilon)$   $\ell_2$ -subspace embedding for the column space of an  $N \times k$  matrix  $A$  is a matrix  $S$  for which for all  $x \in \mathbb{R}^k$ ,  $\|SAx\|_2^2 = (1 \pm \epsilon)\|Ax\|_2^2$ . Equivalently,  $\|I - U^\top S^\top S U\|_2 \leq \epsilon$ , where  $U$  is an orthonormal basis for the column space of  $A$ .*

Combining Fast-JL sketching matrix [AC06, DMM06, Tro11, DMIMW12, LDFU13, PSW17] with a classical  $\epsilon$ -net argument [Woo14] gives subspace embedding,

**Lemma 3.3** (Fast subspace embedding [LDFU13, Woo14]). *Given a matrix  $A \in \mathbb{R}^{N \times k}$  with  $N = \text{poly}(k)$ , then we can compute a  $S \in \mathbb{R}^{k \cdot \text{poly}(\log(k/\delta))/\epsilon^2 \times k}$  that gives a subspace embedding of  $A$  with probability  $1 - \delta$ , i.e., with probability  $1 - \delta$ , we have :*

$$\|SAx\|_2 = (1 \pm \epsilon)\|Ax\|_2$$

holds for any  $x \in \mathbb{R}^n$ ,  $\|x\|_2 = 1$ . Moreover,  $SA$  can be computed in  $O(Nk \cdot \text{poly} \log k)$  time.

## 4 Our Algorithm

Our main algorithm is shown in Algorithm 1. We have the following convergence result of our algorithm.

**Theorem 4.1.** *Suppose the width of a ReLU neural network satisfies*

$$m = \Omega(\max\{\lambda^{-4}n^4, \lambda^{-2}n^2d \log(16n/\delta)\}),$$

*then with probability  $1 - \delta$  over the random initialization of neural network and the randomness of the training algorithm, our algorithm (procedure FASTERTWOLAYER in Algorithm 1) achieves*

$$\|f_{t+1} - y\|_2 \leq \frac{1}{2}\|f_t - y\|_2.$$

*The computation cost in each iteration is  $\tilde{O}(mnd + n^3)$ , and the running time for reducing the training loss to  $\epsilon$  is  $\tilde{O}((mnd + n^3) \log(1/\epsilon))$ . Using fast matrix-multiplication, the total running time can be further reduced to  $\tilde{O}((mnd + n^\omega) \log(1/\epsilon))$ .*

The main difference between [CGH<sup>+</sup>19, ZMG19] and our algorithm is that we perform an *approximate Newton update* (see line 6). The crucial observation here is that the Newton method is robust to small loss, thus it suffices to present a fine approximation. This observation is well-known in the convex optimization but unclear to the non-convex (but overparameterized) neural network setting. Another crucial observation is that instead of directly approximating the Gram matrix, it is suffices to approximate  $(J_t J_t^\top)^{-1} g_t = G_t^{-1} g_t$ . Intuitively, this follows from

$$J_t^\top g_t \approx J_t (J_t J_t^\top)^{-1} (f_t - y) = (J_t^\top J_t)^\dagger J_t (f_t - y),$$

where  $(J_t^\top J_t)^\dagger$  denotes the pseudo-inverse of  $J_t^\top J_t$  and the last term is exactly the Newton update. This observation allows us to formulate the problem a regression problem (see Eq. (6)), on which we can introduce techniques from *randomize linear algebra* and develop fast algorithm that solves it in near linear time.

---

**Algorithm 1** Faster algorithm for two-layer neural network

---

```
1: procedure FASTERTWOLAYER() ▷ Theorem 4.1
2:    $W_0$  is a random Gaussian matrix ▷  $W_0 \in \mathbb{R}^{md}$ 
3:   while  $t < T$  do
4:     Compute the Jacobian matrix  $J_t$  ▷  $J_t \in \mathbb{R}^{n \times md}$ 
5:     Find an  $\epsilon_0$  approximate solution using Algorithm 2 ▷  $\epsilon_0 \in (0, \frac{1}{6}\sqrt{\lambda/n}]$ 


$$\min_{g_t} \|J_t J_t^\top g_t - (f_t - y)\|_2 \tag{6}$$


6:     Update  $W_{t+1} \leftarrow W_t - J_t^\top g_t$ 
7:      $t \leftarrow t + 1$ 
8:   end while
9: end procedure
```

---

---

**Algorithm 2** Fast regression

---

```
1: procedure FASTREGRESSION( $A, \epsilon$ ) ▷ Lemma 4.2
2:    $A \in \mathbb{R}^{N \times k}$  is a full rank matrix,  $\epsilon \in (0, 1/2)$  is the desired precision
3:   Compute a subspace embedding  $SA$  ▷  $S \in \mathbb{R}^{k \text{poly}(\log k) \times k}$ 
4:   Compute  $R$  such that  $SAR$  orthonormal columns via QR decomposition ▷  $R \in \mathbb{R}^{k \times k}$ 
5:    $z_0 \leftarrow \vec{0} \in \mathbb{R}^k$ 
6:   while  $\|A^\top AR z_t - y\|_2 \geq \epsilon$  do
7:      $z_{t+1} \leftarrow z_t - (R^\top A^\top AR)^\top (R^\top A^\top AR z_t - R^\top y)$ 
8:   end while
9:   return  $R z_t$ 
10: end procedure
```

---

#### 4.1 Fast regression solver

The core component of our algorithm is a fast regression solver (shown in Algorithm 2). The regression solver provides an approximate solution to  $\min_x \|A^\top Ax - y\|$  where  $A \in \mathbb{R}^{N \times k}$  ( $N \gg k$ ). We perform preconditioning on the matrix of  $A^\top A$  (line 3 – 4) and use gradient descent to derive an approximation solution (line 6 – 8).

**Lemma 4.2.** *Let  $N = \Omega(k \text{poly}(\log k))$ . Given a matrix  $A \in \mathbb{R}^{N \times k}$ , let  $\kappa$  denote the condition number of  $A$ <sup>5</sup>, consider the following regression problem*

$$\min_{x \in \mathbb{R}^k} \|A^\top Ax - y\|_2. \tag{7}$$

*Using procedure FASTREGRESSION (in Algorithm 2), with probability  $1 - \delta$ , we can compute an  $\epsilon$ -approximate solution  $x'$  satisfying*

$$\|A^\top Ax' - y\|_2 \leq \epsilon \|y\|_2$$

*in  $\tilde{O}(Nk \log(\kappa/\epsilon) + k^3)$  time.*

---

<sup>5</sup> $\kappa = \sigma_{\max}(A)/\sigma_{\min}(A)$

**Speedup in Convex Optimization** It should come as no surprise that our techniques can help accelerating a broad class of solvers in *convex optimization* problems as well. In Section D of the appendix, we elaborate on this application, and in particular show how our technique improves the runtime of the “Newton-Sketch” algorithm of [PW17].

## 5 Conclusion and Open Problems

Our work provides a computationally-efficient (near-linear time) second-order algorithm for training sufficiently overparametrized two-layer neural network, overcoming the drawbacks of traditional first-order gradient algorithms. Our main technical contribution is developing a *faster regression solver* which uses linear sketching for fast preconditioning (in time independent of the network width). As such, our work demonstrates that the toolbox of randomized linear algebra can substantially reduce the computational cost of second-order methods in *non-convex optimization*, and not just in the convex setting for which it was originally developed (e.g., [PW17, Woo14, CLS19, JSWZ20, JKL<sup>+</sup>20, JLSW20, LSZ19]).

Finally, we remark that, while the running time of our algorithm is  $\tilde{O}(Mn + n^3)$  (or  $O(Mn + n^\omega)$  using FMM), it is no longer (near) linear for networks with parameters  $M \leq n^2$  (resp.  $M \lesssim n^{\omega-1}$ ). While it is widely believed that  $\omega = 2$  [CKSU05], FMM algorithms are impractical at present, and it would therefore be very interesting to improve the extra additive term from  $n^3$  to  $n^{2+o(1)}$  (which seems best possible for dense  $n \times n$  matrices), or even to  $n^{3-\epsilon}$  using a practically viable algorithm. Faster preconditioners seem key to this avenue.

## A Appendix

**Organization** The Appendix is organized as follows. Section A contains notations and some basic facts. In Section B we present the fast regression solver. In Section C we prove our main result for two-layer ReLU networks. Finally, in Section D we show that our optimization framework can obtain acceleration in classic convex optimization setting, improve over [PW17].

### A.1 Notation

For a vector  $x \in \mathbb{R}^n$ , we use  $\|x\|_2$  to denote the  $\ell_2$  norm, i.e.,  $\|x\|_2 = (\sum_{i=1}^n x_i^2)^{1/2}$ . We use  $\|x\|_1$  to denote its  $\ell_1$  norm,  $\|x\|_\infty$  to denote its  $\ell_\infty$  norm. For a matrix  $A$ , we use  $\|A\|$  to denote its spectral norm, i.e.,  $\|A\| = \max_{\|x\|_2=1} \|Ax\|_2$ . We use  $\|A\|_F$  to denote the Frobenius norm, i.e.,  $\|A\|_F = (\sum_{i=1}^m \sum_{j=1}^n A_{i,j}^2)^{1/2}$ . We use  $A^\top$  to denote the transpose of matrix  $A$ . We use  $\sigma_{\min}(A)$  to denote the minimum singular value of  $A$ , i.e.,  $\sigma_{\min} = \min_{\|x\|_2=1} \|Ax\|_2$ . We define  $\sigma_{\max}$  to be the maximum singular value and we have  $\sigma_{\max}(A) = \|A\|$ . We use  $\kappa(A)$  to denote the condition number of  $A$ , i.e.,  $\kappa(A) = \sigma_{\max}(A)/\sigma_{\min}(A)$ . We write  $x = y \pm \epsilon$  if  $x \in [y - \epsilon, y + \epsilon]$ . For a positive semidefinite (PSD) matrix  $A$ , we sometimes use  $\lambda_{\min}(A)$  (resp.  $\lambda_{\max}(A)$ ) to denote the minimum (resp. maximum) eigenvalue of  $A$ .

### A.2 Probability Tools

**Lemma A.1** (Chernoff bound [Che52]). *Let  $X = \sum_{i=1}^n X_i$ , where  $X_i = 1$  with probability  $p_i$  and  $X_i = 0$  with probability  $1 - p_i$ , and all  $X_i$  are independent. Let  $\mu = \mathbb{E}[X] = \sum_{i=1}^n p_i$ . Then*

1.  $\Pr[X \geq (1 + \delta)\mu] \leq \exp(-\delta^2\mu/3), \forall \delta > 0$  ;
2.  $\Pr[X \leq (1 - \delta)\mu] \leq \exp(-\delta^2\mu/2), \forall 0 < \delta < 1$ .

**Lemma A.2** (Hoeffding bound [Hoe63]). *Let  $X_1, \dots, X_n$  denote  $n$  independent bounded variables in  $[a_i, b_i]$ . Let  $X = \sum_{i=1}^n X_i$ , then we have*

$$\Pr[|X - \mathbb{E}[X]| \geq t] \leq 2 \exp\left(-\frac{2t^2}{\sum_{i=1}^n (b_i - a_i)^2}\right).$$

**Lemma A.3** (folklore). *Let  $X \sim \mathcal{N}(0, \sigma^2)$ , that is, the probability density function of  $X$  is given by  $\phi(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{x^2}{2\sigma^2}}$ . Then*

$$\Pr[|X| \leq t] \leq \frac{4}{5} \frac{t}{\sigma}.$$

### A.3 Basic Facts

**Fact A.4.** *For any two matrices  $A, B$ ,  $\kappa(B) \leq \kappa(AB)\kappa(A)$ .*

*Proof.* We know for any  $\|x\|_2 = 1$ ,

$$\sigma_{\min}(A)\|Bx\|_2 \leq \|ABx\|_2 \leq \sigma_{\max}(AB)\|x\|_2 = \sigma_{\max}(AB).$$

Hence we have  $\sigma_{\max}(B) \leq \sigma_{\max}(AB)/\sigma_{\min}(A)$ . Similarly, we have

$$\sigma_{\max}(A)\|Bx\|_2 \geq \|ABx\|_2 \geq \sigma_{\min}(AB)\|x\|_2 = \sigma_{\min}(AB)$$

i.e.,  $\sigma_{\min}(B) \geq \sigma_{\min}(AB)/\sigma_{\max}(A)$ . Thus we conclude

$$\kappa(B) \leq \kappa(AB)\kappa(A).$$

□

## B Fast regression solver

**Lemma B.1** (Formal version of Lemma 4.2). *Given a matrix  $A \in \mathbb{R}^{N \times k}$  ( $N \geq k \text{poly}(\log k)$ ), let  $\kappa$  denote the condition number of  $A$ <sup>6</sup>, consider the following regression problem*

$$\min_{x \in \mathbb{R}^k} \|A^\top Ax - y\|_2. \quad (8)$$

*We can compute an  $\epsilon$ -approximate solution  $x'$  satisfying*

$$\|A^\top Ax' - y\|_2 \leq \epsilon \|y\|_2$$

*in  $\tilde{O}(Nk \log(\kappa/\epsilon) + k^3)$  time.*

*Proof.* Using lemma 3.3, let  $S \in \mathbb{R}^{k \text{poly}(\log k/\delta)/\epsilon_0^2 \times N}$  be a subspace embedding of  $A$ , with probability  $1 - \delta$ , the following holds for any  $x \in \mathbb{R}^k$

$$\|SAx\|_2 = (1 \pm \epsilon_0) \|Ax\|_2.$$

Suppose  $R \in \mathbb{R}^{k \times k}$  is computed so that  $SAR$  has orthonormal columns, e.g., via QR decomposition. We use  $R$  as a preconditioner for matrix  $A$ . Formally, for any  $x \in \mathbb{R}^n$  satisfying  $\|x\|_2 = 1$ , we have

$$\|ARx\|_2 = (1 \pm \epsilon_0) \|SARx\|_2 = (1 \pm \epsilon_0). \quad (9)$$

Hence, we know for any  $\|x\|_2 = 1$ ,

$$(1 - \epsilon_0)^2 \leq \|R^\top A^\top ARx\|_2 \leq (1 + \epsilon_0)^2.$$

We choose  $\epsilon_0 = 0.1$ , and consider the regression problem

$$\min_{z \in \mathbb{R}^n} \|R^\top A^\top ARz - R^\top y\|_2. \quad (10)$$

By lemma B.2, using gradient descent, after  $t = \log(1/\epsilon)$  iterations, we can find  $z_t$  satisfying

$$\|R^\top A^\top AR(z_t - z^*)\|_2 \leq \epsilon \|R^\top A^\top AR(z_0 - z^*)\|_2, \quad (11)$$

where  $z^* = (R^\top A^\top AR)^{-1} R^\top y$  is the optimal solution to Eq. (10). We are going to show that  $x_t = Rz_t$  is an  $2\kappa\epsilon$ -approximate solution to the original regression problem (8), i.e.,

$$\|A^\top Ax_t - y\|_2 \leq \kappa\epsilon \|y\|_2$$

Plugging  $z_0 = 0$  into Eq. (11), we get

$$\|R^\top A^\top Ax_t - R^\top y\|_2 \leq \epsilon \|R^\top y\|_2 \leq \epsilon \cdot \sigma_{\max}(R^\top) \|y\|_2 \quad (12)$$

On the other hand, we have

$$\|R^\top A^\top Ax_t - R^\top y\|_2 = \|R^\top (A^\top Ax_t - y)\|_2 \geq \sigma_{\min}(R^\top) \|A^\top Ax_t - y\|_2. \quad (13)$$

---

<sup>6</sup> $\kappa = \sigma_{\max}(A)/\sigma_{\min}(A)$

Putting it all together, we have

$$\|A^\top Ax_t - y\|_2 \leq \epsilon \kappa(R^\top) \|y\|_2 = \epsilon \kappa(R) \|y\|_2 \leq \epsilon \kappa(AR) \kappa(A) \|y\|_2 \leq 2\epsilon \kappa(A) \|y\|_2$$

where the first step follows from Eq. (12) (13), the second step follows from  $R$  is a square matrix and thus  $\kappa(R) = \kappa(R^\top)$ , the third step follows from Fact A.4 and the last step follows from Eq. (9).

For the running time, the preconditioning time is  $\tilde{O}(Nk + k^3)$ , the number of iteration for gradient descent is  $\log(\kappa/\epsilon)$ , the running time per iteration is  $\tilde{O}(Nk)$ , thus the total running time is

$$\tilde{O}(Nk \log(\kappa/\epsilon) + k^3).$$

□

**Lemma B.2.** *Consider the regression problem*

$$\min_x \|Bx - y\|_2^2.$$

Suppose  $B$  is a PSD matrix with  $\frac{3}{4} \leq \|Bx\|_2 \leq \frac{5}{4}$  holds for all  $\|x\|_2 = 1$ . Using gradient descent, after  $t$  iterations, we obtain

$$\|B(x_t - x^*)\|_2 \leq c^t \|B(x_0 - x^*)\|_2$$

for some constant  $c \in (0, 0.9]$ .

*Proof.* The gradient at time  $t$  is  $B^\top(Bx_t - y)$  and  $x_{t+1} = x_t - B^\top(Bx_t - y)$ , thus we have

$$\begin{aligned} \|Bx_{t+1} - Bx^*\|_2 &= \|B(x_t - B^\top(Bx_t - y)) - Bx^*\|_2 \\ &= \|B(x_t - x^*) - BB^\top Bx_t + BB^\top Bx^*\|_2 \\ &= \|(I - BB^\top)B(x_t - x^*)\|_2 \\ &\leq \|I - BB^\top\| \cdot \|B(x_t - x^*)\|_2 \\ &\leq \frac{9}{16} \|B(x_t - x^*)\|_2 \end{aligned}$$

The second step follows from  $B^\top Bx^* = B^\top y$ . The last step follows from the eigenvalue of  $BB^\top$  belongs to  $[\frac{9}{16}, \frac{25}{16}]$  by our assumption. Thus we complete the proof. □

## C Our Algorithm

We dedicate to prove the following result in this section, which is essentially Theorem 4.1.

**Theorem C.1** (Formal version of Theorem 4.1). *Suppose the width of the neural network satisfies  $m = \Omega(\max\{\lambda^{-4}n^4, \lambda^{-2}n^2d \log(16n/\delta)\})$ , then with probability  $1 - \delta$  over the random initialization of neural network and the randomness of the algorithm, our algorithm achieves*

$$\|f_{t+1} - y\|_2 \leq \frac{1}{2} \|f_t - y\|_2.$$

The computation cost in each iteration is  $\tilde{O}(mnd + n^3)$ , and the running time for reducing the training loss to  $\epsilon$  is  $\tilde{O}((mnd + n^3) \log(1/\epsilon))$ . Using fast matrix multiplication, the running time is  $\tilde{O}((mnd + n^\omega) \log(1/\epsilon))$ .

The follow lemmas are standard in literature.

**Lemma C.2** (Bounds on initialization, Lemma 2 in [CGH<sup>+</sup>19]). *Suppose  $m = \Omega(d \log(n/\delta))$ , then with probability  $1 - \delta$ , we have the following*

- $f(W, x_i) = O(1)$ , for  $i \in [n]$ .
- $\|J_{W_0, x_i}\|_F = O(1)$ , for  $i \in [n]$ .

**Lemma C.3** (Bounds on the least eigenvalue at intialization, Lemma 3 in [CGH<sup>+</sup>19]). *Suppose  $m = \Omega(\lambda^{-2} n^2 \log(n/\delta))$ , then with probability at least  $1 - \delta$ , we have*

$$\lambda_{\min}(G_0) \geq \frac{3}{4}\lambda.$$

When weights do not change very much, we have

**Lemma C.4.** *Suppose  $R \geq 1$  and  $m = \tilde{\Omega}(n^2 R^2)$ . With probability at least  $1 - \delta$  over the random initialization of  $W_0$ , the following holds for any set of weights  $w_1, \dots, w_m \in \mathbb{R}^d$  satisfying  $\max_{r \in [m]} \|w_r - w_r(0)\|_2 \leq R/\sqrt{m}$ ,*

- $\|W - W_0\| = O(R)$ ,
- $\|J_{W, x_i} - J_{W_0, x_i}\|_2 = \tilde{O}(R^{1/2}/m^{1/4})$  and  $\|J_W - J_{W_0}\|_F = \tilde{O}(n^{1/2} R^{1/2}/m^{1/4})$ ,
- $\|J_W\|_F = O(\sqrt{n})$ ,

*Proof.* (1) The first claim follows from

$$\|W - W_0\| \leq \|W - W_0\|_F = \left( \sum_{r=1}^m \|w_r - w_r(0)\|_2^2 \right)^{1/2} \leq \sqrt{m} \cdot R/\sqrt{m} = R.$$

(2) For the second claim, we have for any  $i \in [n]$

$$\begin{aligned} \|J_{W, x_i} - J_{W_0, x_i}\|_2^2 &= \frac{1}{m} \sum_{r=1}^m a_r^2 \cdot \|x_r\|_2^2 \cdot |\mathbf{1}_{\langle w_r, x_i \rangle \geq 0} - \mathbf{1}_{\langle w_r(0), x_i \rangle \geq 0}|^2 \\ &= \frac{1}{m} \sum_{r=1}^m |\mathbf{1}_{\langle w_r, x_i \rangle \geq 0} - \mathbf{1}_{\langle w_r(0), x_i \rangle \geq 0}|. \end{aligned} \tag{14}$$

The second equality follows from  $a_r \in \{-1, 1\}$ ,  $\|x_i\|_2 = 1$  and

$$s_{i,r} := |\mathbf{1}_{\langle w_r, x_i \rangle \geq 0} - \mathbf{1}_{\langle w_r(0), x_i \rangle \geq 0}| \in \{0, 1\}. \tag{15}$$

We define the event  $A_{i,r}$  as

$$A_{i,r} = \{\exists \tilde{w} : \|\tilde{w} - w_r(0)\| \leq R/\sqrt{m}, \quad \mathbf{1}_{\langle \tilde{w}, x_i \rangle \geq 0} \neq \mathbf{1}_{\langle w_r(0), x_i \rangle \geq 0}\}.$$

It is easy to see  $A_{i,r}$  happens if and only if  $w_r(0)^\top x_i \in [-R/\sqrt{m}, R/\sqrt{m}]$ . By the anticoncentration of Gaussian (see Lemma A.3), we have  $\mathbb{E}[s_{i,r}] = \Pr[A_{i,r}] \leq \frac{4}{5} R/\sqrt{m}$ . Thus we have

$$\begin{aligned} \Pr \left[ \sum_{i=1}^m s_{i,r} \geq (t + 4/5) R\sqrt{m} \right] &\leq \Pr \left[ \sum_{i=1}^m (s_{i,r} - \mathbb{E}[s_{i,r}]) \geq t R\sqrt{m} \right] \\ &\leq 2 \exp \left( -\frac{2t^2 R^2 m}{m} \right) \\ &= 2 \exp(-t^2 R^2) \\ &\leq 2 \exp(-t^2). \end{aligned} \tag{16}$$

holds for any  $t > 0$ . The second inequality comes from the Hoeffding bound (see Lemma A.2), the last inequality comes from  $R > 1$ . Taking  $t = 2 \log(n/\delta)$  and using union bound over  $i$ , with probability  $1 - \delta$ , we have

$$\|J_{W,x_i} - J_{W_0,x_i}\|_2^2 = \frac{1}{m} \sum_{r=1}^m s_{i,r} \leq \frac{1}{m} \cdot 2 \log(n/\delta) R \sqrt{m} = \tilde{O}(R/\sqrt{m})$$

holds for all  $i \in [n]$ . The first equality comes from Eq. (14) and Eq. (15), the second inequality comes from Eq. (16). Thus we conclude with

$$\|J_{W,x_i} - J_{W_0,x_i}\|_2 = \tilde{O}(R^{1/2}/m^{1/4}) \quad \text{and} \quad \|J_W - J_{W_0}\|_F = \tilde{O}(n^{1/2} R^{1/2}/m^{1/4}).$$

(3) The thrid claim follows from

$$\|J_W\|_F \leq \|J_{W_0}\|_F + \|J_W - J_{W_0}\|_F \leq O(\sqrt{n}) + \tilde{O}(n^{1/2} R^{1/2}/m^{1/4}) = O(\sqrt{n}).$$

The second inequality follows from  $m = \tilde{\Omega}(R^2 n^2)$ . □

**Lemma C.5** (Bounds on the least eigenvalue during optimization, Lemma 4.2 in [SY19]). *Suppose  $m = \Omega(n^2 R^2 \log(n/\delta))$ , with probability at least  $1 - \delta$ , the following holds for any set of weights  $w_1, \dots, w_m \in \mathbb{R}^d$  satisfying  $\max_{r \in [m]} \|w_r - w_r(0)\|_2 \leq R/\sqrt{m}$ ,*

$$\|G_W - G_{W_0}\|_F \leq \lambda/2.$$

We now begin the proof of Theorem C.1

*Proof of Theorem C.1.* We use induction to prove the following two claims recursively. We take  $R \approx n/\lambda$  in the proof.

1.  $\|w_r(t) - w_r(0)\|_2 \leq R/\sqrt{m}$  holds for any  $r \in [m]$  and  $t \geq 0$ .
2.  $\|f_t - y\|_2 \leq \frac{1}{2} \|f_{t-1} - y\|_2$  holds for any  $t \geq 1$ .

Suppose the above two claims hold up to  $t$ , we prove they continue to hold for time  $t + 1$ . The second claim is more delicate, we are going to prove it first and we define

$$J_{t,t+1} = \int_0^1 J((1-s)W_t + sW_{t+1}) ds.$$

Hence, we have

$$\begin{aligned} & \|f_{t+1} - y\|_2 \\ &= \|f_t - y + (f_{t+1} - f_t)\|_2 \\ &= \|f_t - y + J_{t,t+1}(W_{t+1} - W_t)\|_2 \\ &= \|f_t - y - J_{t,t+1} J_t^\top g_t\|_2 \\ &= \|f_t - y - J_t J_t^\top g_t + J_t J_t^\top g_t - J_{t,t+1} J_t^\top g_t\|_2 \\ &\leq \|f_t - y - J_t J_t^\top g_t\|_2 + \|(J_t - J_{t,t+1}) J_t^\top g_t\|_2 \\ &\leq \|f_t - y - J_t J_t^\top g_t\|_2 + \|(J_t - J_{t,t+1}) J_t^\top g^*\|_2 + \|(J_t - J_{t,t+1}) J_t^\top (g_t - g^*)\|_2, \end{aligned} \tag{17}$$

where we denote  $g^* = (J_t J_t^\top)^{-1}(f_t - y)$  to be the optimal solution to Eq. (6). The second step follows from the definition of  $J_{t,t+1}$  and simple calculus. The third step follows from the updating rule of the algorithm.

For the first term of Eq. (17), we have

$$\|J_t J_t^\top g_t - (f_t - y)\|_2 \leq \frac{1}{6} \|f_t - y\|_2, \quad (18)$$

since  $g_t$  is an  $\epsilon_0$  ( $\epsilon_0 \leq \frac{1}{6}$ ) approximate solution to regression problem (6).

For the second term in Eq. (17), we have

$$\begin{aligned} \|(J_t - J_{t,t+1})J_t^\top g^*\|_2 &\leq \|(J_t - J_{t,t+1})\| \cdot \|J_t^\top g^*\|_2 \\ &= \|(J_t - J_{t,t+1})\| \cdot \|J_t^\top (J_t J_t^\top)^{-1}(f_t - y)\|_2 \\ &\leq \|(J_t - J_{t,t+1})\| \cdot \|J_t^\top (J_t J_t^\top)^{-1}\| \cdot \|f_t - y\|_2. \end{aligned} \quad (19)$$

We bound these term separately. First,

$$\begin{aligned} \|J_t - J_{t,t+1}\| &\leq \int_0^1 \|J((1-s)W_t + sW_{t+1}) - J(W_t)\| ds \\ &\leq \int_0^1 (\|J((1-s)W_t + sW_{t+1}) - J(W_0)\| + \|J(W_0) - J(W_t)\|) ds \\ &\leq \tilde{O}(R^{1/2}n^{1/2}/m^{1/4}). \end{aligned} \quad (20)$$

The third step follows from the second claim in Lemma C.4 and the fact that

$$\begin{aligned} \|(1-s)w_r(t) + sw_r(t+1) - w_0\|_2 &\leq (1-s)\|w_r(t) - w_r(0)\|_2 + s\|w_r(t+1) - w_r(0)\|_2 \\ &\leq R/\sqrt{m}. \end{aligned}$$

Furthermore, we have

$$\|J_t^\top (J_t J_t^\top)^{-1}\| = \frac{1}{\sigma_{\min}(J_t^\top)} \leq \sqrt{2/\lambda} \quad (21)$$

The second inequality follows from  $\sigma_{\min}(J_t) = \sqrt{\lambda_{\min}(J_t^\top J_t)} \geq \sqrt{\lambda/2}$  (see Lemma C.5).

Combining Eq. (19), (20) and (21), we have

$$\|(J_t - J_{t,t+1})J_t^\top g^*\|_2 \leq \tilde{O}(R^{1/2}\lambda^{-1}n^{1/2}/m^{1/4})\|f_t - y\|_2 \leq \frac{1}{6}\|f_t - y\|_2, \quad (22)$$

since  $m = \tilde{\Omega}(\lambda^{-4}n^4)$ .

For the third term in Eq. (17), we have

$$\|(J_t - J_{t,t+1})J_t^\top (g_t - g^*)\|_2 \leq \|J_t - J_{t,t+1}\| \cdot \|J_t^\top\| \cdot \|g_t - g^*\|_2. \quad (23)$$

Moreover, one has

$$\begin{aligned} \frac{\lambda}{2}\|g_t - g^*\|_2 &\leq \lambda_{\min}(J_t J_t^\top)\|g_t - g^*\|_2 \\ &\leq \|J_t J_t^\top g_t - J_t J_t^\top g^*\|_2 \\ &= \|J_t J_t^\top g_t - (f_t - y)\|_2 \\ &\leq \sqrt{\lambda/n} \cdot \|f_t - y\|_2. \end{aligned} \quad (24)$$

The first step comes from  $\lambda_{\min}(J_t J_t^\top) = \lambda_{\min}(G_t) \geq \lambda/2$  (see Lemma C.4) and the last step comes from  $g_t$  is an  $\epsilon_0$  approximate solution to Eq. (6). The fourth step follows from Eq. (24) and the fact that  $\|(J_t J_t^\top)^{-1}\| \leq 2/\lambda$ . The last step follows from  $g_t$  is an  $\epsilon_0$  ( $\epsilon_0 \leq \sqrt{\lambda/n}$ ) approximate solution to the regression (6).

Consequently, we have

$$\begin{aligned}
\|(J_t - J_{t,t+1})J_t^\top(g_t - g^*)\|_2 &\leq \|J_t - J_{t,t+1}\| \cdot \|J_t^\top\| \cdot \|g_t - g^*\|_2 \\
&\leq \tilde{O}(R^{1/2}n^{1/2}/m^{1/4}) \cdot \sqrt{n} \cdot \frac{2}{\sqrt{n\lambda}} \cdot \|f_t - y\|_2 \\
&= \tilde{O}(R^{1/2}\lambda^{-1/2}n^{1/2}/m^{1/4}) \cdot \|f_t - y\|_2 \\
&\leq \frac{1}{6}\|f_t - y\|_2
\end{aligned} \tag{25}$$

The second step follows from Eq. (20) and (24) and the fact that  $\|J_t\| \leq O(\sqrt{n})$  (see Lemma C.4). The last step follows from the  $m \geq \Omega(n^4\lambda^{-4})$ . Combining Eq. (17), (18), (22), and (25), we have proved the second claim, i.e.,

$$\|f_{t+1} - y\|_2 \leq \frac{1}{2}\|f_t - y\|_2. \tag{26}$$

It remains to show that  $W_t$  does not move far away from  $W_0$ . First, we have

$$\begin{aligned}
\|g_t\|_2 &\leq \|g^*\|_2 + \|g_t - g^*\|_2 \\
&\leq \|(J_t J_t^\top)^{-1}(f_t - y)\|_2 + \|g_t - g^*\|_2 \\
&\leq \|(J_t J_t^\top)^{-1}\| \cdot \|f_t - y\|_2 + \|g_t - g^*\|_2 \\
&\leq \frac{2}{\lambda} \cdot \|f_t - y\|_2 + \frac{2}{\sqrt{n\lambda}} \cdot \|f_t - y\|_2 \\
&\lesssim \frac{1}{\lambda} \cdot \|f_t - y\|_2
\end{aligned} \tag{27}$$

where the third step follows from Eq. (24) and the last step follows from the obvious fact that  $1/\sqrt{n\lambda} \leq 1/\lambda$ .

Hence, for any  $r \in [m]$  and  $0 \leq k \leq t$ , if we use  $g_{k,i}$  to denote the  $i^{th}$  indice of  $g_k$ , then we have

$$\begin{aligned}
\|w_r(k+1) - w_r(k)\|_2 &= \left\| \sum_{i=1}^n \frac{1}{\sqrt{m}} a_r x_r^\top \mathbf{1}_{\langle w_r(t), x_r \rangle \geq 0} g_{k,i} \right\|_2 \\
&\leq \frac{1}{\sqrt{m}} \sum_{i=1}^n |g_{k,i}| \\
&\leq \frac{\sqrt{n}}{\sqrt{m}} \|g_k\|_2 \\
&\lesssim \frac{\sqrt{n}}{\sqrt{m}} \cdot \frac{1}{2^k \lambda} \|f_0 - y\|_2 \\
&\lesssim \frac{n}{\sqrt{m\lambda}} \cdot \frac{1}{2^k}
\end{aligned}$$

The first step follows from the updating rule, the second step follows from triangle inequalities and the fact that  $a_r = \pm 1$ ,  $\|x_r\|_2 = 1$ . The third step comes from Cauchy-Schwartz inequality, and

the fourth step comes from Eq. (26) and Eq. (27). The last inequality comes from the fact that  $\|f_0 - y\|_2 \leq O(\sqrt{n})$  (see Lemma C.2). Consequently, we have

$$\|w_r(t+1) - w_r(0)\|_2 \leq \sum_{k=0}^t \|w_r(k+1) - w_r(k)\|_2 \lesssim \sum_{k=0}^t \frac{n}{\sqrt{m}\lambda} \cdot \frac{1}{2^k} \lesssim \frac{R}{\sqrt{m}}.$$

Thus we also finish the proof of the first claim.

It remains to give an analysis on the running time of our algorithm. In each iteration, besides evaluating function value and doing backpropagation, which generally takes  $O(mnd)$  time, we also need to solve the regression problem in (6), which takes  $\tilde{O}(mnd \log(\kappa(J_t J_t^\top)/\epsilon_0) + n^3)$  time by Lemma B.1. From Lemma C.4, we know  $\|J_t J_t^\top\| = \|G_t\| \leq O(n)$  and  $\lambda_{\min}(J_t J_t^\top) = \lambda_{\min}(G_t) \geq O(\lambda)$ . Moreover, we only need to set  $\epsilon_0 = \min\{\sqrt{\lambda/n}, 1/6\}$ . Thus the total computation cost in each iteration is  $\tilde{O}(mnd + n^3)$ , and the total running time to reduce the training loss below  $\epsilon$  is  $\tilde{O}((mnd + n^3) \log(1/\epsilon))$ .  $\square$

## D Application: Convex Optimization

We apply our technique to convex optimization problem. We follow the problem formulation in [PW17] and consider the problem

$$\min_x f(x)$$

where  $f$  is  $\gamma$ -strongly convex,  $\beta$ -smooth and its Hessian matrix  $\nabla^2 f(x)$  is  $L$  Lipschitz continuous,

**Definition D.1** ( $\gamma$ -strongly convex). *The function  $f$  is  $\gamma$ -strongly convex if*

$$f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle + \frac{\gamma}{2} \|y - x\|_2^2.$$

**Definition D.2** ( $\beta$ -smooth). *The function  $f$  is  $\beta$ -strongly convex if*

$$f(y) \leq f(x) + \langle \nabla f(x), y - x \rangle + \frac{\beta}{2} \|y - x\|_2^2.$$

**Definition D.3** ( $L$  Lipschitz continuous Hessian). *The Hessian matrix of function  $f$  is  $L$  Lipschitz continuous is*

$$\|\nabla^2 f(x) - \nabla^2 f(y)\| \leq L \|x - y\|_2.$$

As in [PW17], we further assume we have access to

$$\text{the square root of Hessian} := \nabla^2 f(x)^{\frac{1}{2}} \in \mathbb{R}^{m \times n}$$

with  $m \geq n \text{poly}(\log n)$ .

There are many natural and interesting examples that are valid for this assumption. For instance, suppose the objective function has the form of  $f(x) = g(Ax)$  where  $A \in \mathbb{R}^{n \times d}$  and the function  $g : \mathbb{R}^n \rightarrow \mathbb{R}$  has the separable form  $g(Ax) = \sum_{i=1}^n g_i(\langle a_i, x \rangle)$ , then the square root of Hessian is given by

$$\nabla^2 f(x)^{\frac{1}{2}} = D_g(x)A \in \mathbb{R}^{n \times d},$$

where  $D_g(x) \in \mathbb{R}^{n \times n}$  is a diagonal matrix such that the  $i, i$ -th of  $D_g(x)$  is  $\sqrt{g_i''(\langle a_i, x \rangle)}$ .

For more examples, we refer interested reader to Section 3.3 in [PW17]

Naive implementation of Newton method needs to compute

$$\nabla^2 f(x) = (\nabla^2 f(x)^{\frac{1}{2}})^\top \cdot (\nabla^2 f(x)^{\frac{1}{2}})$$

and it costs  $O(nd^2)$  time. The original analysis of NEWTONSKETCH in [PW17] takes  $\tilde{O}(nd + d^3)$ , but it requires  $n \geq d\kappa^2$ , where  $\kappa$  is the condition number defined as  $\kappa = \beta/\gamma$ . There are many follow up work [XYR<sup>+</sup>16, YLZ17, BBN19] intending to get rid of the extra dependence on the condition number  $\kappa$ . We present an alternative approach and improve the running time to  $\tilde{O}((n \log(\kappa) + d^2)d \log(1/\epsilon))$  by incorporating the “fast regression solver” introduced in this paper.

---

**Algorithm 3** Fast Newton Update

---

```

1: procedure FASTNEWTONUPDATE( $f, x_0$ ) ▷ Theorem D.4
2:    $\triangleright x_0$  is an initial point that is satisfying  $\|x_0 - x^*\|_2 \leq O(\gamma/L)$ 
3:    $t \leftarrow 1$ 
4:   while  $t < T$  do
5:     Compute  $\nabla^2 f(x_t)^{\frac{1}{2}} \in \mathbb{R}^{m \times n}$  and  $\nabla f(x) \in \mathbb{R}^n$ .
6:     Find an  $1/(4\kappa)$  approximate solution  $g_t \in \mathbb{R}^n$  to the regression problem
           
$$\min_{g_t \in \mathbb{R}^n} \|(\nabla^2 f(x_t)^{\frac{1}{2}})^\top \cdot (\nabla^2 f(x_t)^{\frac{1}{2}}) \cdot g_t - \nabla f(x_t)\|_2 \quad (28)$$

7:      $x_{t+1} \leftarrow x_t - g_t$ 
8:      $t \leftarrow t + 1$ 
9:   end while
10:  return  $x_T$ 
11: end procedure

```

---

Our algorithm is shown in Algorithm 3. Formally, we have

**Theorem D.4.** *Suppose function  $f$  is  $\gamma$ -strongly convex,  $\beta$ -smooth and its Hessian is  $L$  Lipschitz continuous. Given an initialization point  $x_0$  satisfying  $\|x_0 - x^*\|_2 \leq \gamma/(2L)$ , there is an algorithm (procedure FASTNEWTONUPDATE in Algorithm 3) achieves*

$$\|x_{t+1} - x^*\|_2 \leq \frac{1}{4}\|x_t - x^*\|_2 + \frac{L}{\gamma}\|x_t - x^*\|_2^2, \quad (29)$$

Consequently, in order to find an  $\epsilon$  approximate optimal solution, the running time is

$$\tilde{O}((nd \log(\kappa) + d^3) \log(1/\epsilon)).$$

Using fast matrix multiplication, the running time can be further reduced to  $\tilde{O}((nd \log(\kappa) + d^\omega) \log(1/\epsilon))$ .

*Proof.* We first analyze the correctness, and then give an analysis on the running time. Denote

$$\tilde{g}_t = \nabla^2 f(x_t)^{-1} \nabla f(x_t) = \left( (\nabla^2 f(x_t)^{\frac{1}{2}})^\top \cdot (\nabla^2 f(x_t)^{\frac{1}{2}}) \right)^{-1} \nabla f(x_t). \quad (30)$$

We have

$$\begin{aligned}
& \|\nabla^2 f(x_t)(x_{t+1} - x^*)\|_2 \\
&= \|\nabla^2 f(x_t)(x_t - x^* + x_{t+1} - x_t)\|_2 \\
&= \|\nabla^2 f(x_t)(x_t - x^*) - \nabla^2 f(x_t)g_t\|_2 \\
&= \|\nabla^2 f(x_t)(x_t - x^*) - \nabla^2 f(x_t)\tilde{g}_t + \nabla^2 f(x_t)\tilde{g}_t - \nabla^2 f(x_t)g_t\|_2 \\
&\leq \|\nabla^2 f(x_t)(x_t - x^*) - \nabla^2 f(x_t)\tilde{g}_t\|_2 + \|\nabla^2 f(x_t)\tilde{g}_t - \nabla^2 f(x_t)g_t\|_2
\end{aligned} \tag{31}$$

For the first term

$$\begin{aligned}
& \|\nabla^2 f(x_t)(x_t - x^*) - \nabla^2 f(x_t)\tilde{g}_t\|_2 \\
&= \|\nabla^2 f(x_t)(x_t - x^*) - \nabla f(x_t)\|_2 \\
&= \left\| \nabla^2 f(x_t)(x_t - x^*) - \int_{s=0}^1 \nabla^2 f(x^* + s(x_t - x^*))(x_t - x^*) \mathbf{d}s \right\|_2 \\
&= \left\| \int_{s=0}^1 (\nabla^2 f(x_t) - \nabla^2 f(x^* + s(x_t - x^*))) (x_t - x^*) \mathbf{d}s \right\|_2 \\
&\leq \int_{s=0}^1 \|\nabla^2 f(x_t) - \nabla^2 f(x^* + s(x_t - x^*))\| \mathbf{d}s \cdot \|x_t - x^*\|_2 \\
&\leq \int_{s=0}^1 L(1-s) \|x_t - x^*\|_2 \mathbf{d}s \cdot \|x_t - x^*\|_2 \\
&\leq L \|x_t - x^*\|_2^2.
\end{aligned} \tag{32}$$

The first step follows from the definition of  $\tilde{g}_t$  in Eq. (30), the second step follows from  $\nabla f(x^*) = 0$ . If the Hessian is  $L$  Lipschitz continuous, we have For the second term

$$\begin{aligned}
\|\nabla^2 f(x_t)g_t - \nabla^2 f(x_t)\tilde{g}_t\|_2 &= \|\nabla^2 f(x_t)g_t - \nabla f(x_t)\|_2 \\
&\leq \frac{1}{4\kappa} \|\nabla f(x_t)\|_2 \\
&= \frac{1}{4\kappa} \cdot \beta \|x_t - x^*\|_2 \\
&= \frac{\gamma}{4} \|x_t - x^*\|_2.
\end{aligned} \tag{33}$$

The first step follows from Eq. (30), the second step holds since  $g_t$  is an  $1/(4\kappa)$  approximate solution to Eq. (28). The third step follows from the smoothness of  $f$ . Consequently, we have

$$\begin{aligned}
\|x_{t+1} - x^*\| &\leq \frac{1}{\gamma} \|\nabla^2 f(x_t)(x_{t+1} - x_t)\|_2 \\
&\leq \frac{1}{\gamma} (L \|x_t - x^*\|_2^2 + \frac{\gamma}{4} \|x_t - x^*\|_2) \\
&\leq \frac{1}{4} \|x_t - x^*\|_2 + \frac{L}{\gamma} \|x_t - x^*\|_2^2.
\end{aligned}$$

The first step follows from the convexity of  $f$ . The second step follows from Eq. (31), (32), and (33). Thus we prove the correctness of Eq. (29). Since we know  $\kappa(\nabla^2 f(x_t)^{\frac{1}{2}}) = \sqrt{\kappa}$ , the running time per iteration is  $\tilde{O}(nd \log(\kappa) + d^3)$  by Lemma B.1. Thus we conclude the proof.  $\square$

## References

- [ABH17] Naman Agarwal, Brian Bullins, and Elad Hazan. Second-order stochastic optimization for machine learning in linear time. *The Journal of Machine Learning Research*, 18(1):4148–4187, 2017.
- [AC06] Nir Ailon and Bernard Chazelle. Approximate nearest neighbors and the fast johnson-lindenstrauss transform. In *Proceedings of the thirty-eighth annual ACM symposium on Theory of computing (STOC)*, pages 557–563, 2006.
- [ADH<sup>+</sup>19a] Sanjeev Arora, Simon Du, Wei Hu, Zhiyuan Li, and Ruosong Wang. Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks. In *International Conference on Machine Learning (ICML)*, pages 322–332. <https://arxiv.org/pdf/1901.08584.pdf>, 2019.
- [ADH<sup>+</sup>19b] Sanjeev Arora, Simon S Du, Wei Hu, Zhiyuan Li, Russ R Salakhutdinov, and Ruosong Wang. On exact computation with an infinitely wide neural net. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 8139–8148. <https://arxiv.org/pdf/1904.11955.pdf>, 2019.
- [AGK<sup>+</sup>20] Rohan Anil, Vineet Gupta, Tomer Koren, Kevin Regan, and Yoram Singer. Second order optimization made practical. *arXiv preprint arXiv:2002.09018*, 2020.
- [AKM<sup>+</sup>17] Haim Avron, Michael Kapralov, Cameron Musco, Christopher Musco, Ameya Velingker, and Amir Zandieh. Random fourier features for kernel ridge regression: Approximation bounds and statistical guarantees. In *ICML*. <https://arxiv.org/pdf/1804.09893.pdf>, 2017.
- [AKM<sup>+</sup>19] Haim Avron, Michael Kapralov, Cameron Musco, Christopher Musco, Ameya Velingker, and Amir Zandieh. A universal sampling method for reconstructing signals with simple fourier transforms. In *STOC*. <https://arxiv.org/pdf/1812.08723.pdf>, 2019.
- [ALS<sup>+</sup>18] Alexandr Andoni, Chengyu Lin, Ying Sheng, Peilin Zhong, and Ruiqi Zhong. Subspace embedding and linear regression with orlicz norm. In *ICML*, 2018.
- [AMT10] Haim Avron, Petar Maymounkov, and Sivan Toledo. Blendenpik: Supercharging lapack’s least-squares solver. *SIAM Journal on Scientific Computing*, 32(3):1217–1236, 2010.
- [AZLL19] Zeyuan Allen-Zhu, Yuanzhi Li, and Yingyu Liang. Learning and generalization in overparameterized neural networks, going beyond two layers. In *Advances in neural information processing systems (NeurIPS)*, pages 6155–6166, 2019.
- [AZLS19a] Zeyuan Allen-Zhu, Yuanzhi Li, and Zhao Song. A convergence theory for deep learning via over-parameterization. In *ICML*. <https://arxiv.org/pdf/1811.03962>, 2019.
- [AZLS19b] Zeyuan Allen-Zhu, Yuanzhi Li, and Zhao Song. On the convergence rate of training recurrent neural networks. In *NeurIPS*. <https://arxiv.org/pdf/1810.12065>, 2019.
- [BBN19] Raghu Bollapragada, Richard H Byrd, and Jorge Nocedal. Exact and inexact subsampled newton methods for optimization. *IMA Journal of Numerical Analysis*, 39(2):545–578, 2019.
- [BCW19] Ainesh Bakshi, Nadiia Chepurko, and David P Woodruff. Robust and sample optimal algorithms for psd low-rank approximation. In *arXiv preprint*. <https://arxiv.org/pdf/1912.04177.pdf>, 2019.
- [BELM20] Sébastien Bubeck, Ronen Eldan, Yin Tat Lee, and Dan Mikulincer. Network size and weights size for memorization with two-layers neural networks. In *arXiv preprint*. <https://arxiv.org/pdf/2006.02855.pdf>, 2020.
- [BJW19] Ainesh Bakshi, Rajesh Jayaram, and David P Woodruff. Learning two layer rectified neural networks in polynomial time. In *COLT*. <https://arxiv.org/pdf/1811.01885.pdf>, 2019.

- [BLC88] Sue Becker and Yann Le Cun. Improving the convergence of back-propagation learning with second order methods. In *Proceedings of the 1988 connectionist models summer school*, pages 29–37, 1988.
- [BLH18] Alberto Bernacchia, Máté Lengyel, and Guillaume Hennequin. Exact natural gradient in deep linear networks and its application to the nonlinear case. In *Advances in Neural Information Processing Systems (NIPS)*, pages 5941–5950, 2018.
- [BRB17] Aleksandar Botev, Hippolyt Ritter, and David Barber. Practical gauss-newton optimisation for deep learning. In *International Conference on Machine Learning (ICML)*, pages 557–565, 2017.
- [BW14] Christos Boutsidis and David P Woodruff. Optimal cur matrix decompositions. In *Proceedings of the 46th Annual ACM Symposium on Theory of Computing (STOC)*, pages 353–362. ACM, <https://arxiv.org/pdf/1405.7910>, 2014.
- [BW18] Ainesh Bakshi and David Woodruff. Sublinear time low-rank approximation of distance matrices. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 3782–3792. <https://arxiv.org/pdf/1809.06986.pdf>, 2018.
- [BWZ16] Christos Boutsidis, David P Woodruff, and Peilin Zhong. Optimal principal component analysis in distributed and streaming models. In *Proceedings of the forty-eighth annual ACM symposium on Theory of Computing (STOC)*, pages 236–249, 2016.
- [BWZ20] Frank Ban, David P. Woodruff, and Richard Zhang. Regularized weighted low rank approximation. In *NeurIPS*. <https://arxiv.org/pdf/1911.06958.pdf>, 2020.
- [CG19] Yuan Cao and Quanquan Gu. Generalization bounds of stochastic gradient descent for wide and deep neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 10835–10845, 2019.
- [CGH<sup>+</sup>19] Tianle Cai, Ruiqi Gao, Jikai Hou, Siyu Chen, Dong Wang, Di He, Zhihua Zhang, and Liwei Wang. A gram-gauss-newton method learning overparameterized deep neural networks for regression problems. In *arXiv preprint*. <https://arXiv.org/pdf/1905.11675>, 2019.
- [Che52] Herman Chernoff. A measure of asymptotic efficiency for tests of a hypothesis based on the sum of observations. *The Annals of Mathematical Statistics*, pages 493–507, 1952.
- [CKSU05] Henry Cohn, Robert Kleinberg, Balazs Szegedy, and Christopher Umans. Group-theoretic algorithms for matrix multiplication. In *46th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 379–388. IEEE, 2005.
- [CLS19] Michael B Cohen, Yin Tat Lee, and Zhao Song. Solving linear programs in the current matrix multiplication time. In *STOC*. <https://arxiv.org/pdf/1810.07896>, 2019.
- [CW13] Kenneth L. Clarkson and David P. Woodruff. Low rank approximation and regression in input sparsity time. In *Symposium on Theory of Computing Conference (STOC)*, pages 81–90. <https://arxiv.org/pdf/1207.6365>, 2013.
- [Dan20] Amit Daniely. Memorizing gaussians with no over-parameterizaion via gradient decent on neural networks. In *arXiv preprint*. <https://arxiv.org/pdf/2003.12895.pdf>, 2020.
- [DHS11] John Duchi, Elad Hazan, and Yoram Singer. Adaptive subgradient methods for online learning and stochastic optimization. *Journal of machine learning research (JMLR)*, 12(Jul):2121–2159, 2011.
- [DJS<sup>+</sup>19] Huaian Diao, Rajesh Jayaram, Zhao Song, Wen Sun, and David Woodruff. Optimal sketching for kronecker product regression and low rank approximation. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 4739–4750. <https://arxiv.org/pdf/1909.13384.pdf>, 2019.

- [DMIMW12] Petros Drineas, Malik Magdon-Ismail, Michael W Mahoney, and David P Woodruff. Fast approximation of matrix coherence and statistical leverage. *Journal of Machine Learning Research*, 13(Dec):3475–3506, 2012.
- [DMM06] Petros Drineas, Michael W Mahoney, and Shan Muthukrishnan. Sampling algorithms for l2 regression and applications. In *Proceedings of the seventeenth annual ACM-SIAM symposium on Discrete algorithm*, pages 1127–1136. Society for Industrial and Applied Mathematics, 2006.
- [DZPS19] Simon S Du, Xiyu Zhai, Barnabas Poczos, and Aarti Singh. Gradient descent provably optimizes over-parameterized neural networks. In *ICLR*. <https://arxiv.org/pdf/1810.02054.pdf>, 2019.
- [GKS18] Vineet Gupta, Tomer Koren, and Yoram Singer. Shampoo: Preconditioned stochastic tensor optimization. In *International Conference on Machine Learning (ICML)*, pages 1842–1850, 2018.
- [GLB<sup>+</sup>18] Thomas George, César Laurent, Xavier Bouthillier, Nicolas Ballas, and Pascal Vincent. Fast approximate natural gradient descent in a kronecker factored eigenbasis. In *Advances in Neural Information Processing Systems (NIPS)*, pages 9550–9560, 2018.
- [GM16] Roger Grosse and James Martens. A kronecker-factored approximate fisher matrix for convolution layers. In *International Conference on Machine Learning (ICML)*, pages 573–582, 2016.
- [GU18] François Le Gall and Florent Urrutia. Improved rectangular matrix multiplication using powers of the coppersmith-winograd tensor. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1029–1046. <https://arxiv.org/pdf/1708.05622.pdf>, 2018.
- [Hoe63] Wassily Hoeffding. Probability inequalities for sums of bounded random variables. *Journal of the American Statistical Association*, 58(301):13–30, 1963.
- [HRS16] Moritz Hardt, Benjamin Recht, and Yoram Singer. Train faster, generalize better: Stability of stochastic gradient descent. In *ICML*. <https://arxiv.org/pdf/1509.01240.pdf>, 2016.
- [JGH18] Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems (NIPS)*, pages 8571–8580, 2018.
- [JKL<sup>+</sup>20] Haotian Jiang, Tarun Kathuria, Yin Tat Lee, Swati Padmanabhan, and Zhao Song. A faster interior point method for semidefinite programming. In *Manuscript*, 2020.
- [JLSW20] Haotian Jiang, Yin Tat Lee, Zhao Song, and Sam Chiu-wai Wong. An improved cutting plane method for convex optimization, convex-concave games and its applications. In *STOC*. <https://arxiv.org/pdf/2004.04250.pdf>, 2020.
- [JSWZ20] Shunhua Jiang, Zhao Song, Omri Weinstein, and Hengjie Zhang. Faster dynamic matrix inverse for faster lps. In *arXiv preprint*. <https://arxiv.org/pdf/2004.07470.pdf>, 2020.
- [JT20] Ziwei Ji and Matus Telgarsky. Polylogarithmic width suffices for gradient descent to achieve arbitrarily small test error with shallow relu networks. In *ICLR*. <https://arxiv.org/pdf/1909.12292.pdf>, 2020.
- [KB15] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *ICLR*. <https://arxiv.org/pdf/1412.6980.pdf>, 2015.
- [KMP10] Ioannis Koutis, Gary L Miller, and Richard Peng. Approaching optimality for solving sdd linear systems. In *2010 IEEE 51st Annual Symposium on Foundations of Computer Science*. IEEE, 2010.
- [KMP11] Ioannis Koutis, Gary L Miller, and Richard Peng. A nearly-m log n time solver for sdd linear systems. In *2011 IEEE 52nd Annual Symposium on Foundations of Computer Science*, pages 590–598. IEEE, 2011.

- [KOSZ13] Jonathan A Kelner, Lorenzo Orecchia, Aaron Sidford, and Zeyuan Allen Zhu. A simple, combinatorial algorithm for solving sdd systems in nearly-linear time. In *Proceedings of the forty-fifth annual ACM symposium on Theory of computing (STOC)*, pages 911–920. ACM, <https://arxiv.org/pdf/1301.6628.pdf>, 2013.
- [KS16] Rasmus Kyng and Sushant Sachdeva. Approximate gaussian elimination for laplacians-fast, sparse, and simple. In *IEEE 57th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 573–582. <https://arxiv.org/pdf/1605.02353.pdf>, 2016.
- [LDFU13] Yichao Lu, Paramveer Dhillon, Dean P Foster, and Lyle Ungar. Faster ridge regression via the subsampled randomized hadamard transform. In *Advances in neural information processing systems*, pages 369–377, 2013.
- [LG14] François Le Gall. Powers of tensors and fast matrix multiplication. In *Proceedings of the 39th international symposium on symbolic and algebraic computation (ISSAC)*, pages 296–303. ACM, 2014.
- [LJH<sup>+</sup>19] Liyuan Liu, Haoming Jiang, Pengcheng He, Weizhu Chen, Xiaodong Liu, Jianfeng Gao, and Jiawei Han. On the variance of the adaptive learning rate and beyond. In *arXiv preprint*. <https://arxiv.org/pdf/1908.03265.pdf>, 2019.
- [LL18] Yuanzhi Li and Yingyu Liang. Learning overparameterized neural networks via stochastic gradient descent on structured data. In *NeurIPS*. <https://arxiv.org/pdf/1808.01204.pdf>, 2018.
- [LN17] Kasper Green Larsen and Jelani Nelson. Optimality of the johnson-lindenstrauss lemma. In *IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 633–638. <https://arxiv.org/pdf/1609.02094.pdf>, 2017.
- [LPPW20] Hang Liao, Barak A. Pearlmutter, Vamsi K. Potluru, and David P. Woodruff. Automatic differentiation of sketched regression. In *AISTATS*, 2020.
- [LSZ19] Yin Tat Lee, Zhao Song, and Qiuyu Zhang. Solving empirical risk minimization in the current matrix multiplication time. In *COLT*. <https://arxiv.org/pdf/1905.04447>, 2019.
- [LWYZ20] Yi Li, Ruosong Wang, Lin Yang, and Hanrui Zhang. Nearly linear row sampling algorithm for quantile regression. In *ICML*. <https://arxiv.org/pdf/2006.08397.pdf>, 2020.
- [LY17] Yuanzhi Li and Yang Yuan. Convergence analysis of two-layer neural networks with ReLU activation. In *Advances in neural information processing systems (NIPS)*, pages 597–607. <https://arxiv.org/pdf/1705.09886.pdf>, 2017.
- [Mar10] James Martens. Deep learning via hessian-free optimization. In *ICML*, volume 27, pages 735–742, 2010.
- [MBJ18] James Martens, Jimmy Ba, and Matthew Johnson. Kronecker-factored curvature approximations for recurrent neural networks. 2018.
- [MG15] James Martens and Roger Grosse. Optimizing neural networks with kronecker-factored approximate curvature. In *International conference on machine learning (ICML)*, pages 2408–2417, 2015.
- [MNJ16] Philipp Moritz, Robert Nishihara, and Michael Jordan. A linearly-convergent stochastic l-bfgs algorithm. In *Artificial Intelligence and Statistics*, pages 249–258, 2016.
- [NN13] Jelani Nelson and Huy L Nguyễn. Osnap: Faster numerical linear algebra algorithms via sparser subspace embeddings. In *2013 IEEE 54th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 117–126. IEEE, <https://arxiv.org/pdf/1211.1002>, 2013.
- [OS19] Samet Oymak and Mahdi Soltanolkotabi. Towards moderate overparameterization: global convergence guarantees for training shallow neural networks. *IEEE Journal on Selected Areas in Information Theory*, 2019.

- [PSW17] Eric Price, Zhao Song, and David P. Woodruff. Fast regression with an  $\ell_\infty$  guarantee. In *International Colloquium on Automata, Languages, and Programming (ICALP)*. <https://arxiv.org/pdf/1705.10723.pdf>, 2017.
- [PW17] Mert Pilanci and Martin J Wainwright. Newton sketch: A near linear-time optimization algorithm with linear-quadratic convergence. *SIAM Journal on Optimization*, 27(1):205–245, 2017.
- [RSW16] Ilya Razenshteyn, Zhao Song, and David P Woodruff. Weighted low rank approximations with provable guarantees. In *Proceedings of the 48th Annual Symposium on the Theory of Computing (STOC)*, 2016.
- [RT08] Vladimir Rokhlin and Mark Tygert. A fast randomized algorithm for overdetermined linear least-squares regression. *Proceedings of the National Academy of Sciences*, 105(36):13212–13217, 2008.
- [Sar06] Tamás Sarlós. Improved approximation algorithms for large matrices via random projections. In *Proceedings of 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2006.
- [Son19] Zhao Song. *Matrix Theory : Optimization, Concentration and Algorithms*. PhD thesis, The University of Texas at Austin, 2019.
- [ST04] Daniel A. Spielman and Shang-Hua Teng. Nearly-linear time algorithms for graph partitioning, graph sparsification, and solving linear systems. In *Proceedings of the Thirty-sixth Annual ACM Symposium on Theory of Computing (STOC)*, pages 81–90. ACM, 2004.
- [SWY<sup>+</sup>19] Zhao Song, Ruosong Wang, Lin Yang, Hongyang Zhang, and Peilin Zhong. Efficient symmetric norm regression via linear sketching. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 828–838. <https://arxiv.org/pdf/1910.01788.pdf>, 2019.
- [SWZ17] Zhao Song, David P Woodruff, and Peilin Zhong. Low rank approximation with entrywise  $\ell_1$ -norm error. In *Proceedings of the 49th Annual Symposium on the Theory of Computing (STOC)*. ACM, <https://arxiv.org/pdf/1611.00898>, 2017.
- [SWZ19a] Zhao Song, David Woodruff, and Peilin Zhong. Average case column subset selection for entrywise  $\ell_1$ -norm loss. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 10111–10121. <https://arxiv.org/pdf/2004.07986.pdf>, 2019.
- [SWZ19b] Zhao Song, David Woodruff, and Peilin Zhong. Towards a zero-one law for column subset selection. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 6120–6131. <https://arxiv.org/pdf/1811.01442.pdf>, 2019.
- [SWZ19c] Zhao Song, David P Woodruff, and Peilin Zhong. Relative error tensor low rank approximation. In *Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 2772–2789. <https://arxiv.org/pdf/1704.08246.pdf>, 2019.
- [SY19] Zhao Song and Xin Yang. Quadratic suffices for over-parametrization via matrix chernoff bound. In *arXiv preprint*. <https://arxiv.org/pdf/1906.03593.pdf>, 2019.
- [Tro11] Joel A Tropp. Improved analysis of the subsampled randomized hadamard transform. *Advances in Adaptive Data Analysis*, 3(01n02):115–126, 2011.
- [Vai89a] Pravin M Vaidya. A new algorithm for minimizing convex functions over convex sets. In *30th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 338–343. IEEE, 1989.
- [Vai89b] Pravin M Vaidya. Speeding-up linear programming using fast matrix multiplication. In *30th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 332–337. IEEE, 1989.
- [WDW19] Xiaoxia Wu, Simon S Du, and Rachel Ward. Global convergence of adaptive gradient methods for an over-parameterized neural network. In *arXiv preprint*. <https://arxiv.org/pdf/1902.07111.pdf>, 2019.

- [Wil12] Virginia Vassilevska Williams. Multiplying matrices faster than coppersmith-winograd. In *Proceedings of the forty-fourth annual ACM symposium on Theory of computing (STOC)*, pages 887–898. ACM, 2012.
- [WMG<sup>+</sup>17] Yuhuai Wu, Elman Mansimov, Roger B Grosse, Shun Liao, and Jimmy Ba. Scalable trust-region method for deep reinforcement learning using kronecker-factored approximation. In *Advances in neural information processing systems (NIPS)*, pages 5279–5288, 2017.
- [Woo14] David P. Woodruff. Sketching as a tool for numerical linear algebra. *Foundations and Trends in Theoretical Computer Science*, 10(1-2):1–157, 2014.
- [WW19] Ruosong Wang and David P Woodruff. Tight bounds for  $\ell_p$  oblivious subspace embeddings. In *Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1825–1843. SIAM, <https://arxiv.org/pdf/1801.04414.pdf>, 2019.
- [WZ16] David P Woodruff and Peilin Zhong. Distributed low rank approximation of implicit functions of a matrix. In *2016 IEEE 32nd International Conference on Data Engineering (ICDE)*, pages 847–858. IEEE, 2016.
- [WZ20] David P. Woodruff and Amir Zandieh. Near input sparsity time kernel embeddings via adaptive sampling. In *ICML*, 2020.
- [XYR<sup>+</sup>16] Peng Xu, Jiyan Yang, Fred Roosta, Christopher Ré, and Michael W Mahoney. Sub-sampled newton methods with non-uniform sampling. In *Advances in Neural Information Processing Systems (NIPS)*, pages 3000–3008, 2016.
- [YLZ17] Haishan Ye, Luo Luo, and Zhihua Zhang. Approximate newton methods and their local convergence. In *International Conference on Machine Learning (ICML)*, pages 3931–3939, 2017.
- [ZG19] Difan Zou and Quanquan Gu. An improved analysis of training over-parameterized deep neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 2053–2062, 2019.
- [ZMG19] Guodong Zhang, James Martens, and Roger B Grosse. Fast convergence of natural gradient descent for over-parameterized neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 8080–8091, 2019.
- [ZPD<sup>+</sup>20] Yi Zhang, Orestis Plevrakis, Simon S Du, Xingguo Li, Zhao Song, and Sanjeev Arora. Over-parameterized adversarial training: An analysis overcoming the curse of dimensionality. In *arXiv preprint*. <https://arxiv.org/pdf/2002.06668.pdf>, 2020.
- [ZSD17] Kai Zhong, Zhao Song, and Inderjit S Dhillon. Learning non-overlapping convolutional neural networks with multiple kernels. In *arXiv preprint*. <https://arxiv.org/pdf/1711.03440.pdf>, 2017.
- [ZSJ<sup>+</sup>17] Kai Zhong, Zhao Song, Prateek Jain, Peter L. Bartlett, and Inderjit S. Dhillon. Recovery guarantees for one-hidden-layer neural networks. In *ICML*. <https://arxiv.org/pdf/1706.03175.pdf>, 2017.