
Optimization of Graph Neural Networks: Implicit Acceleration by Skip Connections and More Depth

Keyulu Xu^{*1} Mozhi Zhang² Stefanie Jegelka¹ Kenji Kawaguchi^{*3}

Abstract

Graph Neural Networks (GNNs) have been studied through the lens of expressive power and generalization. However, their optimization properties are less well understood. We take the first step towards analyzing GNN training by studying the gradient dynamics of GNNs. First, we analyze linearized GNNs and prove that despite the non-convexity of training, convergence to a global minimum at a linear rate is guaranteed under mild assumptions that we validate on real-world graphs. Second, we study what may affect the GNNs’ training speed. Our results show that the training of GNNs is implicitly accelerated by skip connections, more depth, and/or a good label distribution. Empirical results confirm that our theoretical results for linearized GNNs align with the training behavior of nonlinear GNNs. Our results provide the first theoretical support for the success of GNNs with skip connections in terms of optimization, and suggest that deep GNNs with skip connections would be promising in practice.

1. Introduction

Graph Neural Networks (GNNs) (Gori et al., 2005; Scarselli et al., 2009) are an effective framework for learning with graphs. GNNs learn node representations on a graph by extracting high-level features not only from a node itself but also from a node’s surrounding subgraph. Specifically, the node representations are recursively aggregated and updated using neighbor representations (Merkwirth & Lengauer, 2005; Duvenaud et al., 2015; Defferrard et al., 2016; Kearnes et al., 2016; Gilmer et al., 2017; Hamilton et al., 2017; Velickovic et al., 2018; Liao et al., 2020).

Recently, there has been a surge of interest in studying the

theoretical aspects of GNNs to understand their success and limitations. Existing works have studied GNNs’ expressive power (Keriven & Peyré, 2019; Maron et al., 2019; Chen et al., 2019; Xu et al., 2019; Sato et al., 2019; Loukas, 2020), generalization capability (Scarselli et al., 2018; Du et al., 2019b; Xu et al., 2020; Garg et al., 2020), and extrapolation properties (Xu et al., 2021). However, the understanding of the optimization properties of GNNs has remained limited. For example, researchers working on the fundamental problem of designing more expressive GNNs hope and often empirically observe that more powerful GNNs better fit the training set (Xu et al., 2019; Sato et al., 2020; Vignac et al., 2020). Theoretically, given the non-convexity of GNN training, it is still an open question whether better representational power always translates into smaller training loss. This motivates the more general questions:

*Can gradient descent find a global minimum for GNNs?
What affects the speed of convergence in training?*

In this work, we take an initial step towards answering the questions above by analyzing the trajectory of gradient descent, i.e., *gradient dynamics* or *optimization dynamics*. A complete understanding of the dynamics of GNNs, and deep learning in general, is challenging. Following prior works on gradient dynamics (Saxe et al., 2014; Arora et al., 2019a; Bartlett et al., 2019), we consider the linearized regime, i.e., GNNs with *linear* activation. Despite the linearity, key properties of nonlinear GNNs are present: The objective function is *non-convex* and the dynamics are *nonlinear* (Saxe et al., 2014; Kawaguchi, 2016). Moreover, we observe the learning curves of linear GNNs and ReLU GNNs are surprisingly similar, both converging to nearly zero training loss at the same linear rate (Figure 1). Similarly, prior works report comparable performance in node classification benchmarks even if we remove the non-linearities (Thekumparampil et al., 2018; Wu et al., 2019). Hence, understanding the dynamics of linearized GNNs is a valuable step towards understanding the general GNNs.

Our analysis leads to an affirmative answer to the first question. We establish that gradient descent training of a linearized GNN with squared loss converges to a global minimum at a linear rate. Experiments confirm that the assump-

^{*}Equal contribution ¹Massachusetts Institute of Technology (MIT) ²The University of Maryland ³Harvard University. Correspondence to: Keyulu Xu <keyulu@mit.edu>, Kenji Kawaguchi <kkawaguchi@fas.harvard.edu>.

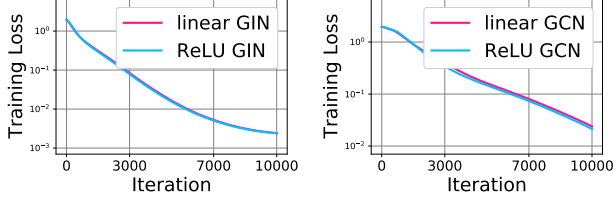


Figure 1. Training curves of linearized GNNs vs. ReLU GNNs on the Cora node classification dataset.

tions of our theoretical results for global convergence hold on real-world datasets. The most significant contribution of our convergence analysis is on multiscale GNNs, i.e., GNN architectures that use *skip connections* to combine graph features at various scales (Xu et al., 2018; Li et al., 2019; Abu-El-Haija et al., 2020; Chen et al., 2020; Li et al., 2020). The skip connections introduce complex interactions among layers, and thus the resulting dynamics are more intricate. To our knowledge, our results are the first convergence results for GNNs with *more than one* hidden layer, with or without skip connections.

We then study what may affect the training speed of GNNs. First, for any fixed depth, GNNs with skip connections train faster. Second, increasing the depth further accelerates the training of GNNs. Third, faster training is obtained when the labels are more correlated with the graph features, i.e., labels contain “signal” instead of “noise”. Overall, experiments for nonlinear GNNs agree with the prediction of our theory for linearized GNNs.

Our results provide the first theoretical justification for the empirical success of multiscale GNNs in terms of optimization, and suggest that deeper GNNs with skip connections may be promising in practice. In the GNN literature, skip connections are initially motivated by the “over-smoothing” problem (Xu et al., 2018): via the recursive neighbor aggregation, node representations of a *deep* GNN on expander-like subgraphs would be mixing features from almost the entire graph, and may thereby “wash out” relevant local information. In this case, shallow GNNs may perform better. Multiscale GNNs with *skip connections* can combine and adapt to the graph features at various scales, i.e., the output of intermediate GNN layers, and such architectures are shown to help with this over-smoothing problem (Xu et al., 2018; Li et al., 2019; 2020; Abu-El-Haija et al., 2020; Chen et al., 2020). However, the properties of multiscale GNNs have mostly been understood at a conceptual level. Xu et al. (2018) relate the learned representations to random walk distributions and Oono & Suzuki (2020) take a boosting view, but they do not consider the optimization dynamics. We give an explanation from the lens of optimization. The training losses of deeper GNNs may be worse due to over-smoothing. In contrast, multiscale GNNs can express any

shallower GNNs and fully exploit the power by converging to a global minimum. Hence, our results suggest that deeper GNNs with skip connections are guaranteed to train faster with smaller training losses.

We present our results on global convergence in Section 3, after introducing relevant background (Section 2). In Section 4, we compare the training speed of GNNs as a function of skip connections, depth, and the label distribution. All proofs are deferred to the Appendix.

2. Preliminaries

2.1. Notation and Background

We begin by introducing our notation. Let $G = (V, E)$ be a graph with n vertices $V = \{v_1, v_2, \dots, v_n\}$. Its adjacency matrix $A \in \mathbb{R}^{n \times n}$ has entries $A_{ij} = 1$ if $(v_i, v_j) \in E$ and 0 otherwise. The degree matrix associated with A is $D = \text{diag}(d_1, d_2, \dots, d_n)$ with $d_i = \sum_{j=1}^n A_{ij}$. For any matrix $M \in \mathbb{R}^{m \times m'}$, we denote its j -th column vector by $M_{*j} \in \mathbb{R}^m$, its i -th row vector by $M_{i*} \in \mathbb{R}^{m'}$, and its largest and smallest (i.e., $\min(m, m')$ -th largest) singular values by $\sigma_{\max}(M)$ and $\sigma_{\min}(M)$, respectively. The data matrix $X \in \mathbb{R}^{m_x \times n}$ has columns X_{*j} corresponding to the feature vector of node v_j , with input dimension m_x .

The task of interest is node classification or regression. Each node $v_i \in V$ has an associated label $y_i \in \mathbb{R}^{m_y}$. In the transductive (semi-supervised) setting, we have access to training labels for only a subset $\mathcal{I} \subseteq [n]$ of nodes on G , and the goal is to predict the labels for the other nodes in $[n] \setminus \mathcal{I}$. Our problem formulation easily extends to the inductive setting by letting $\mathcal{I} = [n]$, and we can use the trained model for prediction on unseen graphs. Hence, we have access to $\bar{n} = |\mathcal{I}| \leq n$ training labels $Y = [y_i]_{i \in \mathcal{I}} \in \mathbb{R}^{m_y \times \bar{n}}$, and we train the GNN using X, Y, G . Additionally, for any $M \in \mathbb{R}^{m \times m'}$, $\mathcal{I}$ may index sub-matrices $M_{*\mathcal{I}} = [M_{*i}]_{i \in \mathcal{I}} \in \mathbb{R}^{m \times \bar{n}}$ (when $m' \geq n$) and $M_{\mathcal{I}*} = [M_{i*}]_{i \in \mathcal{I}} \in \mathbb{R}^{\bar{n} \times m}$ (when $m \geq n$).

Graph Neural Networks (GNNs) use the graph structure and node features to learn representations of nodes (Scarselli et al., 2009). GNNs maintain hidden representations $h_{(l)}^v \in \mathbb{R}^{m_l}$ for each node v , where m_l is the hidden dimension on the l -th layer. We let $X_{(l)} = [h_{(l)}^1, h_{(l)}^2, \dots, h_{(l)}^n] \in \mathbb{R}^{m_l \times n}$, and set $X_{(0)}$ as the input features X . The node hidden representations $X_{(l)}$ are updated by aggregating and transforming the neighbor representations:

$$X_{(l)} = \sigma(B_{(l)}X_{(l-1)}S) \in \mathbb{R}^{m_l \times n}, \quad (1)$$

where σ is a nonlinearity such as ReLU, $B_{(l)} \in \mathbb{R}^{m_l \times m_{l-1}}$ is the weight matrix, and $S \in \mathbb{R}^{n \times n}$ is the GNN aggregation matrix, whose formula depends on the exact variant of GNN. In Graph Isomorphism Networks (GIN) (Xu et al.,

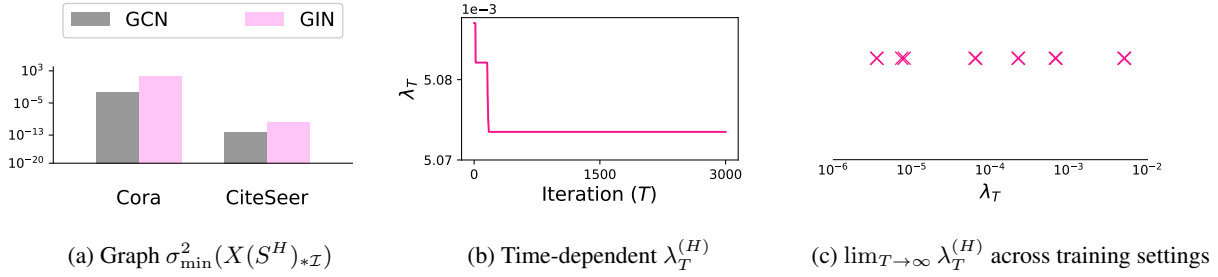


Figure 2. **Empirical validation of assumptions for global convergence of linear GNNs.** Left panel confirms the graph condition $\sigma_{\min}^2(X(S^H)_{*\mathcal{I}}) > 0$ for datasets Cora and CiteSeer, and for models GCN and GIN. Middle panel shows the time-dependent $\lambda_T^{(H)}$ for one training setting (linear GCN on Cora). Each point in right panel is $\lambda_T^{(H)} > 0$ at the last iteration for different training settings.

2019), $S = A + I_n$ is the adjacency matrix of G with self-loop, where $I_n \in \mathbb{R}^{n \times n}$ is an identity matrix. In Graph Convolutional Networks (GCN) (Kipf & Welling, 2017), $S = \hat{D}^{-\frac{1}{2}}(A + I_n)\hat{D}^{-\frac{1}{2}}$ is the normalized adjacency matrix, where $\hat{D}$ is the degree matrix of $A + I_n$.

2.2. Problem Setup

We first formally define linearized GNNs.

Definition 1. (Linear GNN). Given data matrix $X \in \mathbb{R}^{m_x \times n}$, aggregation matrix $S \in \mathbb{R}^{n \times n}$, weight matrices $W \in \mathbb{R}^{m_y \times m_H}$, $B_{(l)} \in \mathbb{R}^{m_l \times m_{l-1}}$, and their collection $B = (B_{(1)}, \dots, B_{(H)})$, a linear GNN with H layers $f(X, W, B) \in \mathbb{R}^{m_y \times n}$ is defined as

$$f(X, W, B) = WX_{(H)}, \quad X_{(l)} = B_{(l)}X_{(l-1)}S. \quad (2)$$

Throughout this paper, we refer multiscale GNNs to the commonly used *Jumping Knowledge Network (JK-Net)* (Xu et al., 2018), which connects the output of all intermediate GNN layers to the final layer with skip connections:

Definition 2. (Multiscale linear GNN). Given data $X \in \mathbb{R}^{m_x \times n}$, aggregation matrix $S \in \mathbb{R}^{n \times n}$, weight matrices $W_{(l)} \in \mathbb{R}^{m_y \times m_l}$, $B_{(l)} \in \mathbb{R}^{m_l \times m_{l-1}}$ with $W = (W_{(0)}, W_{(1)}, \dots, W_{(H)})$, a multiscale linear GNN with H layers $f(X, W, B) \in \mathbb{R}^{m_y \times n}$ is defined as

$$f(X, W, B) = \sum_{l=0}^H W_{(l)}X_{(l)}, \quad (3)$$

$$X_{(l)} = B_{(l)}X_{(l-1)}S. \quad (4)$$

Given a GNN $f(\cdot)$ and a loss function $\ell(\cdot, Y)$, we can train the GNN by minimizing the training loss $L(W, B)$:

$$L(W, B) = \ell(f(X, W, B)_{*\mathcal{I}}, Y), \quad (5)$$

where $f(X, W, B)_{*\mathcal{I}}$ corresponds to the GNN’s predictions on nodes that have training labels and thus incur training

losses. The pair (W, B) represents the trainable weights:

$$L(W, B) = L(W_{(1)}, \dots, W_{(H)}, B_{(1)}, \dots, B_{(H)})$$

For completeness, we define the global minimum of GNNs.

Definition 3. (Global minimum). For any $H \in \mathbb{N}_0$, L_H^* is the global minimum value of the H -layer linear GNN f :

$$L_H^* = \inf_{W, B} \ell(f(X, W, B)_{*\mathcal{I}}, Y). \quad (6)$$

Similarly, we define $L_{1:H}^*$ as the global minimum value of the multiscale linear GNN f with H layers.

We are ready to present our main results on global convergence for linear GNNs and multiscale linear GNNs.

3. Convergence Analysis

In this section, we show that gradient descent training a linear GNN with squared loss, with or without skip connections, converges linearly to a global minimum. Our conditions for global convergence hold on real-world datasets and provably hold under assumptions, e.g., initialization.

In linearized GNNs, the loss $L(W, B)$ is non-convex (and non-invex) despite the linearity. The graph aggregation S creates interaction among the data and poses additional challenges in the analysis. We show a fine-grained analysis of the GNN’s gradient dynamics can overcome these challenges. Following previous works on gradient dynamics (Saxe et al., 2014; Huang & Yau, 2020; Ji & Telgarsky, 2020; Kawaguchi, 2021), we analyze the GNN learning process via the *gradient flow*, i.e., gradient descent with infinitesimal steps: $\forall t \geq 0$, the network weights evolve as

$$\frac{d}{dt}W_t = -\frac{\partial L}{\partial W}(W_t, B_t), \quad \frac{d}{dt}B_t = -\frac{\partial L}{\partial B}(W_t, B_t), \quad (7)$$

where (W_t, B_t) represents the trainable parameters at time t with initialization (W_0, B_0) .

3.1. Linearized GNNs

Theorem 1 states our result on global convergence for linearized GNNs without skip connections.

Theorem 1. *Let f be an H -layer linear GNN and $\ell(q, Y) = \|q - Y\|_F^2$ where $q, Y \in \mathbb{R}^{m_y \times \bar{n}}$. Then, for any $T > 0$,*

$$\begin{aligned} L(W_T, B_T) - L_H^* & \\ \leq (L(W_0, B_0) - L_H^*) e^{-4\lambda_T^{(H)} \sigma_{\min}^2(X(S^H)_{*\mathcal{I}})T}, \end{aligned} \quad (8)$$

where $\lambda_T^{(H)}$ is the smallest eigenvalue $\lambda_T^{(H)} := \inf_{t \in [0, T]} \lambda_{\min}((\bar{B}_t^{(1:H)})^\top \bar{B}_t^{(1:H)})$ and $\bar{B}^{(1:l)} := B_{(l)} B_{(l-1)} \cdots B_{(1)}$ for any $l \in \{0, \dots, H\}$ with $\bar{B}^{(1:0)} := I$.

Proof. (Sketch) We decompose the gradient dynamics into three components: the graph interaction, non-convex factors, and convex factors. We then bound the effects of the graph interaction and non-convex factors through $\sigma_{\min}^2(X(S^H)_{*\mathcal{I}})$ and $\lambda_{\min}((\bar{B}_t^{(1:H)})^\top \bar{B}_t^{(1:H)})$ respectively. The complete proof is in Appendix A.1. $\square$

Theorem 1 implies that convergence to a global minimum at a linear rate is guaranteed if $\sigma_{\min}^2(X(S^H)_{*\mathcal{I}}) > 0$ and $\lambda_T > 0$. The first condition on the product of X and S^H indexed by $\mathcal{I}$ only depends on the node features X and the GNN aggregation matrix S . It is satisfied if $\text{rank}(X(S^H)_{*\mathcal{I}}) = \min(m_x, \bar{n})$, because $\sigma_{\min}(X(S^H)_{*\mathcal{I}})$ is the $\min(m_x, \bar{n})$ -th largest singular value of $X(S^H)_{*\mathcal{I}} \in \mathbb{R}^{m_x \times \bar{n}}$. The second condition $\lambda_T^{(H)} > 0$ is time-dependent and requires a more careful treatment. Linear convergence is implied as long as $\lambda_{\min}((\bar{B}_t^{(1:H)})^\top \bar{B}_t^{(1:H)}) \geq \epsilon > 0$ for all times t before stopping.

Empirical validation of conditions. We verify both the graph condition $\sigma_{\min}^2(X(S^H)_{*\mathcal{I}}) > 0$ and the time-dependent condition $\lambda_T^{(H)} > 0$ for (discretized) $T > 0$. First, on the popular graph datasets, Cora and Citeseer (Sen et al., 2008), and the GNN models, GCN (Kipf & Welling, 2017) and GIN (Xu et al., 2019), we have $\sigma_{\min}^2(X(S^H)_{*\mathcal{I}}) > 0$ (Figure 2a). Second, we train linear GCN and GIN on Cora and Citeseer to plot an example of how the $\lambda_T^{(H)} = \inf_{t \in [0, T]} \lambda_{\min}((\bar{B}_t^{(1:H)})^\top \bar{B}_t^{(1:H)})$ changes with respect to time T (Figure 2b). We further confirm that $\lambda_T^{(H)} > 0$ until convergence, $\lim_{T \rightarrow \infty} \lambda_T^{(H)} > 0$ across different settings, e.g., datasets, depths, models (Figure 2c). Our experiments use the squared loss, random initialization, learning rate 1e-4, and set the hidden dimension to the input dimension (note that Theorem 1 assumes the hidden dimension is at least the input dimension). Further experimental details are in Appendix C. Along with Theorem 1, we conclude that linear GNNs converge linearly

to a global minimum. Empirically, we indeed see both linear and ReLU GNNs converging at the same linear rate to nearly zero training loss in node classification tasks (Figure 1).

Guarantee via initialization. Besides the empirical verification, we theoretically show that a *good initialization* guarantees the time-dependent condition $\lambda_T > 0$ for any $T > 0$. Indeed, like other neural networks, GNNs do not converge to a global optimum with certain initializations: e.g., initializing all weights to zero leads to zero gradients and $\lambda_T^{(H)} = 0$ for all T , and hence no learning. We introduce a notion of *singular margin* and say an initialization is good if it has a positive singular margin. Intuitively, a good initialization starts with an already small loss.

Definition 4. (Singular margin). The initialization (W_0, B_0) is said to have singular margin $\gamma > 0$ with respect to a layer $l \in \{1, \dots, H\}$ if $\sigma_{\min}(B_{(l)} B_{(l-1)} \cdots B_{(1)}) \geq \gamma$ for all (W, B) such that $L(W, B) \leq L(W_0, B_0)$.

Proposition 1 then states that an initialization with positive singular margin γ guarantees $\lambda_T^{(H)} \geq \gamma^2 > 0$ for all T :

Proposition 1. *Let f be a linear GNN with H layers and $\ell(q, Y) = \|q - Y\|_F^2$. If the initialization (W_0, B_0) has singular margin $\gamma > 0$ with respect to the layer H and $m_H \geq m_x$, then $\lambda_T^{(H)} \geq \gamma^2$ for all $T \in [0, \infty)$.*

Proposition 1 follows since $L(W_t, B_t)$ is non-increasing with respect to time t (proof in Appendix A.2).

Relating to previous works, our singular margin is a generalized variant of the deficiency margin of linear feedforward networks (Arora et al., 2019a, Definition 2 and Theorem 1):

Proposition 2. (Informal) *If initialization (W_0, B_0) has deficiency margin $c > 0$, then it has singular margin $\gamma > 0$.*

The formal version of Proposition 2 is in Appendix A.3.

To summarize, Theorem 1 along with Proposition 1 implies that we have a prior guarantee of linear convergence to a global minimum for any graph with $\text{rank}(X(S^H)_{*\mathcal{I}}) = \min(m_x, \bar{n})$ and initialization (W_0, B_0) with singular margin $\gamma > 0$: i.e., for any desired $\epsilon > 0$, we have that $L(W_T, B_T) - L_H^* \leq \epsilon$ for any T such that

$$T \geq \frac{1}{4\gamma^2 \sigma_{\min}^2(X(S^H)_{*\mathcal{I}})} \log \frac{L(A_0, B_0) - L_H^*}{\epsilon}. \quad (9)$$

While the margin condition theoretically guarantees linear convergence, empirically, we have already seen that the convergence conditions of across different training settings for widely used random initialization.

Theorem 1 suggests that the convergence rate depends on a combination of data features X , the GNN architecture and graph structure via S and H , the label distribution and initialization via λ_T . For example, GIN has better such constants than GCN on the Cora dataset with everything else

held equal (Figure 2a). Indeed, in practice, GIN converges faster than GCN on Cora (Figure 1). In general, the computation and comparison of the rates given by Theorem 1 requires computation such as those in Figure 2. In Section 4, we will study an alternative way of comparing the speed of training by directly comparing the gradient dynamics.

3.2. Multiscale Linear GNNs

Without skip connections, the GNNs under linearization still behave like linear feedforward networks with augmented graph features. With skip connections, the dynamics and analysis become much more intricate. The expressive power of multiscale linear GNNs changes significantly as depth increases. Moreover, the skip connections create complex interactions among different layers and graph structures of various scales in the optimization dynamics. Theorem 2 states our convergence results for multiscale linear GNNs in three cases: (i) a general form; (ii) a weaker condition for boundary cases that uses $\lambda_T^{H'}$ instead of $\lambda_T^{1:H}$; (iii) a faster rate if we have monotonic expressive power as depth increases.

Theorem 2. *Let f be a multiscale linear GNN with H layers and $\ell(q, Y) = \|q - Y\|_F^2$ where $q, Y \in \mathbb{R}^{m_y \times \bar{n}}$. Let $\lambda_T^{(1:H)} := \min_{0 \leq l \leq H} \lambda_T^{(l)}$. For any $T > 0$, the following hold:*

(i) (General). *Let $G_H := [X^\top, (XS)^\top, \dots, (XSH)^\top]^\top \in \mathbb{R}^{(H+1)m_x \times n}$. Then*

$$\begin{aligned} L(W_T, B_T) - L_{1:H}^* & \\ \leq (L(W_0, B_0) - L_{1:H}^*) e^{-4\lambda_T^{(1:H)} \sigma_{\min}^2((G_H)_{*\mathcal{I}}) T}. \end{aligned} \quad (10)$$

(ii) (Boundary cases). *For any $H' \in \{0, 1, \dots, H\}$,*

$$\begin{aligned} L(W_T, B_T) - L_{H'}^* & \\ \leq (L(W_0, B_0) - L_{H'}^*) e^{-4\lambda_T^{(H')} \sigma_{\min}^2(X(S^{H'})_{*\mathcal{I}}) T}. \end{aligned} \quad (11)$$

(iii) (Monotonic expressive power). *If there exist $l, l' \in \{0, \dots, H\}$ with $l < l'$ such that $L_l^* \geq L_{l+1}^* \geq \dots \geq L_{l'}^*$ or $L_l^* \leq L_{l+1}^* \leq \dots \leq L_{l'}^*$, then*

$$\begin{aligned} L(W_T, B_T) - L_{l''}^* & \\ \leq (L(W_0, B_0) - L_{l''}^*) e^{-4 \sum_{k=l}^{l'} \lambda_T^{(k)} \sigma_{\min}^2(X(S^k)_{*\mathcal{I}}) T}, \end{aligned} \quad (12)$$

where $l'' = l$ if $L_l^* \geq L_{l+1}^* \geq \dots \geq L_{l'}^*$, and $l'' = l'$ if $L_l^* \leq L_{l+1}^* \leq \dots \leq L_{l'}^*$.

Proof. (Sketch) A key observation in our proof is that the interactions of different scales cancel out to point towards a specific direction in the gradient dynamics induced in a space of the loss value. The complete proof is in Appendix A.4. $\square$

Similar to Theorem 1 for linear GNNs, the most general form (i) of Theorem 2 implies that convergence to the global minimum value of the *entire* multiscale linear GNN $L_{1:H}^*$ at linear rate is guaranteed when $\sigma_{\min}^2((G_H)_{*\mathcal{I}}) > 0$ and $\lambda_T^{(1:H)} > 0$. The graph condition $\sigma_{\min}^2((G_H)_{*\mathcal{I}}) > 0$ is satisfied if $\text{rank}((G_H)_{*\mathcal{I}}) = \min(m_x(H+1), \bar{n})$. The time-dependent condition $\lambda_T^{(1:H)} > 0$ is guaranteed if the initialization (W_0, B_0) has singular margin $\gamma > 0$ with respect to *every* layer (Proposition 3 is proved in Appendix A.5):

Proposition 3. *Let f be a multiscale linear GNN and $\ell(q, Y) = \|q - Y\|_F^2$. If the initialization (W_0, B_0) has singular margin $\gamma > 0$ with respect to every layer $l \in [H]$ and $m_l \geq m_x$ for $l \in [H]$, then $\lambda_T^{(1:H)} \geq \gamma^2$ for all $T \in [0, \infty)$.*

We demonstrate that the conditions of Theorem 2 (i) hold for real-world datasets, suggesting in practice multiscale linear GNNs converge linearly to a global minimum.

Empirical validation of conditions. On datasets Cora and Citeseer and for GNN models GCN and GIN, we confirm that $\sigma_{\min}^2((G_H)_{*\mathcal{I}}) > 0$ (Figure 3a). Moreover, we train multiscale linear GCN and GIN on Cora and Citeseer to plot an example of how the $\lambda_T^{(1:H)}$ changes with respect to time T (Figure 3b), and we confirm that at convergence, $\lambda_T^{(1:H)} > 0$ across different settings (Figure 3c). Experimental details are in Appendix C.

Boundary cases. Because the global minimum value of multiscale linear GNNs $L_{1:H}^*$ can be smaller than that of linear GNNs L_H^* , the conditions in Theorem 2(i) may sometimes be stricter than those of Theorem 1. For example, in Theorem 2(i), we require $\lambda_T^{(1:H)} := \min_{0 \leq l \leq H} \lambda_T^{(l)}$ rather than $\lambda_T^{(H)}$ to be positive. If $\lambda_T^{(l)} = 0$ for some l , then Theorem 2(i) will not guarantee convergence to $L_{1:H}^*$.

Although the boundary cases above did not occur on the tested real-world graphs (Figure 3), for theoretical interest, Theorem 2(ii) guarantees that in such cases, multiscale linear GNNs still converge to a value no worse than the global minimum value of *non-multiscale* linear GNNs. For any intermediate layer H' , assuming $\sigma_{\min}^2(X(S^{H'})_{*\mathcal{I}}) > 0$ and $\lambda_T^{(H')} > 0$, Theorem 2(ii) bounds the loss of the multiscale linear GNN $L(W_T, B_T)$ at convergence by the global minimum value $L_{H'}^*$ of the corresponding linear GNN with H' layers.

Faster rate under monotonic expressive power. Theorem 2(iii) considers a special case that is likely in real graphs: the global minimum value of the non-multiscale linear GNN $L_{H'}^*$ is *monotonic* as H' increases. Then (iii) gives a *faster rate* than (ii) and linear GNNs. For example, if the globally optimal value decreases as linear GNNs get deeper. i.e., $L_0^* \geq L_1^* \geq \dots \geq L_H^*$, or vice versa, $L_0^* \leq L_1^* \leq \dots \leq$

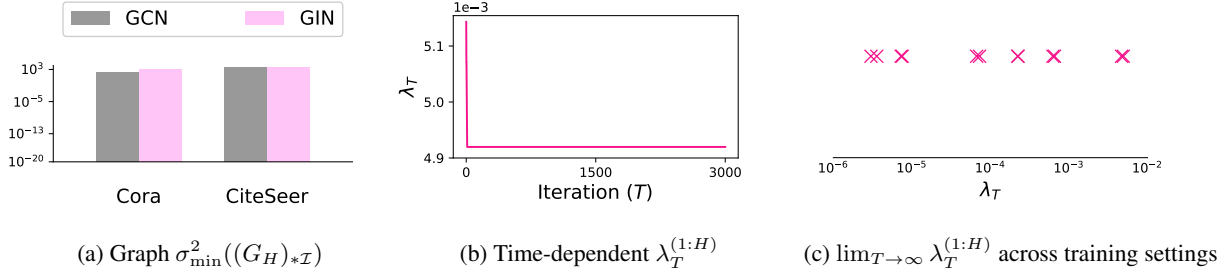


Figure 3. **Empirical validation of assumptions for global convergence of multiscale linear GNNs.** Left panel confirms the graph condition $\sigma_{\min}^2((G_H)_{*\mathcal{I}}) > 0$ for Cora and CiteSeer, and for GCN and GIN. Middle panel shows the time-dependent $\lambda_T^{(1:H)}$ for one training setting (multiscale linear GCN on Cora). Each point in right panel is $\lambda_T^{(1:H)} > 0$ at the last iteration for different training settings.

L_H^* , then Theorem 2 (i) implies that

$$\begin{aligned} L(W_T, B_T) - L_l^* \\ \leq (L(W_0, B_0) - L_l^*)e^{-4 \sum_{k=0}^H \lambda_T^{(k)} \sigma_{\min}^2(X(S^k)_{*\mathcal{I}})T}, \end{aligned} \quad (13)$$

where $l = 0$ if $L_0^* \geq L_1^* \geq \dots \geq L_H^*$, and $l = H$ if $L_0^* \leq L_1^* \leq \dots \leq L_H^*$. Moreover, if the globally optimal value does not change with respect to the depth as $L_{1:H}^* = L_1^* = L_2^* = \dots = L_H^*$, then we have

$$\begin{aligned} L(W_T, B_T) - L_{1:H}^* \\ \leq (L(W_0, B_0) - L_{1:H}^*)e^{-4 \sum_{k=0}^H \lambda_T^{(k)} \sigma_{\min}^2(X(S^k)_{*\mathcal{I}})T}. \end{aligned} \quad (14)$$

We obtain a faster rate for multiscale linear GNNs than for linear GNNs, as $e^{-4 \sum_{k=0}^H \lambda_T^{(k)} \sigma_{\min}^2(X(S^k)_{*\mathcal{I}})T} \leq e^{-4 \lambda_T^{(H)} \sigma_{\min}^2(X(S^H)_{*\mathcal{I}})T}$. Interestingly, unlike linear GNNs, multiscale linear GNNs in this case do not require any condition on initialization to obtain a prior guarantee on global convergence since $e^{-4 \sum_{k=0}^H \lambda_T^{(k)} \sigma_{\min}^2(X(S^k)_{*\mathcal{I}})T} \leq e^{-4 \lambda_T^{(0)} \sigma_{\min}^2(X(S^0)_{*\mathcal{I}})T}$ with $\lambda_T^{(0)} = 1$ and $X(S^0)_{*\mathcal{I}} = X_{*\mathcal{I}}$.

To summarize, we prove global convergence rates for multiscale linear GNNs (Thm. 2(i)) and experimentally validate the conditions. Part (ii) addresses boundary cases where the conditions of Part (i) do not hold. Part (iii) gives faster rates assuming monotonic expressive power with respect to depth. So far, we have shown multiscale linear GNNs converge faster than linear GNNs in the case of (iii). Next, we compare the training speed for more general cases.

4. Implicit Acceleration

In this section, we study how the skip connections, depth of GNN, and label distribution may affect the speed of training for GNNs. Similar to previous works (Arora et al., 2018), we compare the training speed by comparing the per step loss reduction $\frac{d}{dt}L(W_t, B_t)$ for arbitrary differentiable loss functions $\ell(\cdot, Y) : \mathbb{R}^{m_y} \rightarrow \mathbb{R}$. Smaller $\frac{d}{dt}L(W_t, B_t)$ implies faster training. Loss reduction offers a complementary

view to the convergence rates in Section 3, since it is instant and not an upper bound.

We present an analytical form of the loss reduction $\frac{d}{dt}L(W_t, B_t)$ for linear GNNs and multiscale linear GNNs. The comparison of training speed then follows from our formula for $\frac{d}{dt}L(W_t, B_t)$. For better exposition, we first introduce several notations. We let $\bar{B}^{(l':l)} = B_{(l)}B_{(l-1)} \dots B_{(l')}$ for all l' and l where $\bar{B}^{(l':l)} = I$ if $l' > l$. We also define

$$\begin{aligned} J_{(i,l),t} &:= [\bar{B}_t^{(1:i-1)} \otimes (W_{(l),t} \bar{B}_t^{(i+1:l)})^\top], \\ F_{(l),t} &:= [(\bar{B}_t^{(1:l)})^\top \bar{B}_t^{(1:l)} \otimes I_{m_y}] \succeq 0, \\ V_t &:= \frac{\partial L(W_t, B_t)}{\partial \hat{Y}_t}, \end{aligned}$$

where $\hat{Y}_t := f(X, W_t, B_t)_{*\mathcal{I}}$. For any vector $v \in \mathbb{R}^m$ and positive semidefinite matrix $M \in \mathbb{R}^{m \times m}$, we use $\|v\|_M^2 := v^\top M v$.¹ Intuitively, V_t represents the derivative of the loss $L(W_t, B_t)$ with respect to the model output $\hat{Y} = f(X, W_t, B_t)_{*\mathcal{I}}$. $J_{(i,l),t}$ and $F_{(l),t}$ represent matrices that describe how the errors are propagated through the weights of the networks.

Theorem 3, proved in Appendix A.6, gives an analytical formula of loss reduction for linear GNNs and multiscale linear GNNs.

Theorem 3. *For any differentiable loss function $q \mapsto \ell(q, Y)$, the following hold for any $H \geq 0$ and $t \geq 0$:*

(i) (Non-multiscale) For f as in Definition 1:

$$\begin{aligned} \frac{d}{dt}L_1(W_t, B_t) &= -\|\text{vec}[V_t(X(S^H)_{*\mathcal{I}})^\top]\|_{F_{(H),t}}^2 \\ &\quad - \sum_{i=1}^H \|J_{(i,H),t} \text{vec}[V_t(X(S^H)_{*\mathcal{I}})^\top]\|_2^2. \end{aligned} \quad (15)$$

¹We use this Mahalanobis norm notation for conciseness without assuming it to be a norm, since M may be low rank.

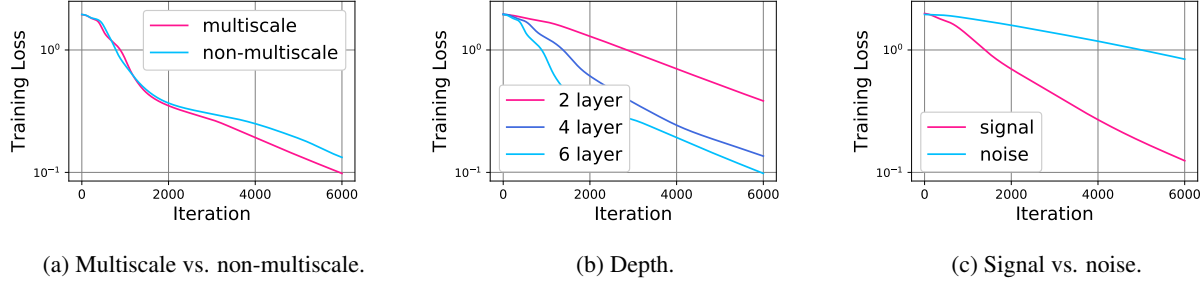


Figure 4. **Comparison of the training speed of GNNs.** Left: Multiscale GNNs train faster than non-multiscale GNNs. Middle: Deeper GNNs train faster. Right: GNNs train faster when the labels have signals instead of random noise. The patterns above hold for both ReLU and linear GNNs. Additional results are in Appendix B.

(ii) (Multiscale) For f as in Definition 2:

$$\begin{aligned} \frac{d}{dt} L_2(W_t, B_t) &= - \sum_{l=0}^H \left\| \text{vec} [V_t(X(S^l)_{*\mathcal{I}})^\top] \right\|_{F(l),t}^2 \\ &\quad - \sum_{i=1}^H \left\| \sum_{l=i}^H J_{(i,l),t} \text{vec} [V_t(X(S^l)_{*\mathcal{I}})^\top] \right\|_2^2. \end{aligned} \quad (16)$$

In what follows, we apply Theorem 3 to predict how different factors affect the training speed of GNNs.

4.1. Acceleration with Skip Connections

We first show that multiscale linear GNNs tend to achieve faster loss reduction $\frac{d}{dt} L_2(W_t, B_t)$ compared to the corresponding linear GNN without skip connections, $\frac{d}{dt} L_1(W_t, B_t)$. It follows from Theorem 3 that

$$\begin{aligned} \frac{d}{dt} L_2(W_t, B_t) - \frac{d}{dt} L_1(W_t, B_t) \\ \leq - \sum_{l=0}^{H-1} \left\| \text{vec} [V_t(X(S^l)_{*\mathcal{I}})^\top] \right\|_{F(l),t}^2, \end{aligned} \quad (17)$$

if $\sum_{i=1}^H (\|a_i\|_2^2 + 2b_i^\top a_i) \geq 0$, where $a_i = \sum_{l=i}^{H-1} J_{(i,l),t} \text{vec}[V_t(X(S^l)_{*\mathcal{I}})^\top]$, and $b_i = J_{(i,H),t} \text{vec}[V_t(X(S^H)_{*\mathcal{I}})^\top]$. The assumption of $\sum_{i=1}^H (\|a_i\|_2^2 + 2b_i^\top a_i) \geq 0$ is satisfied in various ways: for example, it is satisfied if the last layer's term b_i and the other layers' terms a_i are aligned as $b_i^\top a_i \geq 0$, or if the last layer's term b_i is dominated by the other layers' terms a_i as $2\|b_i\|_2 \leq \|a_i\|_2$. Then equation (17) shows that the multiscale linear GNN decreases the loss value with strictly many more negative terms, suggesting faster training.

Empirically, we indeed observe that multiscale GNNs train faster (Figure 4a), both for (nonlinear) ReLU and linear GNNs. We verify this by training multiscale and non-multiscale, ReLU and linear GCNs on the Cora and Citeseer datasets with cross-entropy loss, learning rate 5e-5, and hidden dimension 32. Results are in Appendix B.

4.2. Acceleration with More Depth

Our second finding is that deeper GNNs, with or without skip connections, train faster. For any differentiable loss function $q \mapsto \ell(q, Y)$, Theorem 3 states that the loss of the multiscale linear GNN decreases as

$$\begin{aligned} \frac{d}{dt} L(W_t, B_t) &= - \underbrace{\sum_{l=0}^H \left\| \text{vec} [V_t(X(S^l)_{*\mathcal{I}})^\top] \right\|_{F(l),t}^2}_{\geq 0, \text{ further improvement as depth } H \text{ increases}} \\ &\quad - \underbrace{\sum_{i=1}^H \left\| \sum_{l=i}^H J_{(i,l),t} \text{vec} [V_t(X(S^l)_{*\mathcal{I}})^\top] \right\|_2^2}_{\geq 0, \text{ further improvement as depth } H \text{ increases}}. \end{aligned} \quad (18)$$

In equation (18), we can see that the multiscale linear GNN achieves faster loss reduction as depth H increases. A similar argument applies to non-multiscale linear GNNs.

Empirically too, deeper GNNs train faster (Figure 4b). Again, the acceleration applies to both (nonlinear) ReLU GNNs and linear GNNs. We verify this by training multiscale and non-multiscale, ReLU and linear GCNs with 2, 4, and 6 layers on the Cora and Citeseer datasets with learning rate 5e-5, hidden dimension 32, and cross-entropy loss. Results are in Appendix B.

4.3. Label Distribution: Signal vs. Noise

Finally, we study how the labels affect the training speed. For the loss reduction (15) and (16), we argue that the norm of $V_t(X(S^l)_{*\mathcal{I}})^\top$ tends to be larger for labels Y that are more correlated with the graph features $X(S^l)_{*\mathcal{I}}$, e.g., labels are signals instead of “noise”.

Without loss of generality, we assume Y is normalized, e.g., one-hot labels. Here, $V_t = \frac{\partial L(A_t, B_t)}{\partial \hat{Y}_t}$ is the derivative of the loss with respect to the model output, e.g., $V_t = 2(\hat{Y}_t - Y)$ for squared loss. If the rows of Y are random noise vectors,

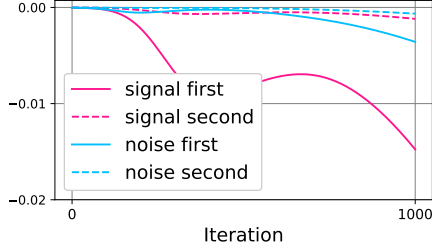


Figure 5. The scale of the first term dominates the second term of the loss reduction $\frac{d}{dt}L(W_t, B_t)$ for linear GNNs trained with the original labels vs. random labels on Cora.

then so are the rows of V_t , and they are expected to get more orthogonal to the columns of $(X(S^l)_{*\mathcal{I}})^\top$ as n increases. In contrast, if the labels Y are highly correlated with the graph features $(X(S^l)_{*\mathcal{I}})^\top$, i.e., the labels have signal, then the norm of $V_t(X(S^l)_{*\mathcal{I}})^\top$ will be larger, implying faster training.

Our argument above focuses on the first term of the loss reduction, $\|V_t(X(S^l)_{*\mathcal{I}})^\top\|_F^2$. We empirically demonstrate that the scale of the second term, $\left\|\sum_{l=i}^H J_{(i,l),t} \text{vec}[V_t(X(S^l)_{*\mathcal{I}})^\top]\right\|_2^2$, is dominated by that of the first term (Figure 5). Thus, we can expect GNNs to train faster with signals than noise.

We train GNNs with the original labels of the dataset and random labels (i.e., selecting a class with uniform probability), respectively. The prediction of our theoretical analysis aligns with practice: training is much slower for random labels (Figure 4c). We verify this for multiscale and non-multiscale, ReLU and linear GCNs on the Cora and Citseer datasets with learning rate $1e-4$, hidden dimension 32, and cross-entropy loss. Results are in Appendix B.

5. Related Work

Theoretical analysis of linearized networks. The theoretical study of neural networks with some linearized components has recently drawn much attention. Tremendous efforts have been made to understand linear *feedforward* networks, in terms of their loss landscape (Kawaguchi, 2016; Hardt & Ma, 2017; Laurent & Brecht, 2018) and optimization dynamics (Saxe et al., 2014; Arora et al., 2019a; Bartlett et al., 2019; Du & Hu, 2019; Zou et al., 2020). Recent works prove global convergence rates for deep linear networks under certain conditions (Bartlett et al., 2019; Du & Hu, 2019; Arora et al., 2019a; Zou et al., 2020). For example, Arora et al. (2019a) assume the data to be whitened. Zou et al. (2020) fix the weights of certain layers during training. Our work is inspired by these works but differs in that our analysis applies to all learnable weights and does not require

these specific assumptions, and we study the more complex GNN architecture with skip connections. GNNs consider the interaction of graph structures via the recursive message passing, but such structured, locally varying interaction is not present in feedforward networks. Furthermore, linear feedforward networks, even with skip connections, have the same expressive power as shallow linear models, a crucial condition in previous proofs (Bartlett et al., 2019; Du & Hu, 2019; Arora et al., 2019a; Zou et al., 2020). In contrast, the expressive power of multiscale linear GNNs can change significantly as depth increases. Accordingly, our proofs significantly differ from previous studies.

Another line of works studies the gradient dynamics of neural networks in the neural tangent kernel (NTK) regime (Jacot et al., 2018; Li & Liang, 2018; Allen-Zhu et al., 2019; Arora et al., 2019b; Chizat et al., 2019; Du et al., 2019a,c; Kawaguchi & Huang, 2019; Nitanda & Suzuki, 2021). With over-parameterization, the NTK remains almost constant during training. Hence, the corresponding neural network is implicitly linearized with respect to random features of the NTK at initialization (Lee et al., 2019; Yehudai & Shamir, 2019; Liu et al., 2020). On the other hand, our work needs to address nonlinear dynamics and changing expressive power.

Learning dynamics and optimization of GNNs. Closely related to our work, Du et al. (2019b); Xu et al. (2021) study the gradient dynamics of GNNs via the Graph NTK but focus on GNNs’ generalization and extrapolation properties. We instead analyze optimization. Only Zhang et al. (2020) also prove global convergence for GNNs, but for the one-hidden-layer case, and they assume a specialized tensor initialization and training algorithms. In contrast, our results work for any finite depth with no assumptions on specialized training. Other works aim to accelerate and stabilize the training of GNNs through normalization techniques (Cai et al., 2020) and importance sampling (Chen et al., 2018a,b; Huang et al., 2018; Chiang et al., 2019; Zou et al., 2019). Our work complements these practical works with a better theoretical understanding of GNN training.

6. Conclusion

This work studies the training properties of GNNs through the lens of optimization dynamics. For linearized GNNs with or without skip connections, despite the non-convex objective, we show that gradient descent training is guaranteed to converge to a global minimum at a linear rate. The conditions for global convergence are validated on real-world graphs. We further find out that skip connections, more depth, and/or a good label distribution implicitly accelerate the training of GNNs. Our results suggest deeper GNNs with skip connections may be promising in practice, and serve as a first foundational step for understanding the optimization of general GNNs.

Acknowledgements

KX and SJ were supported by NSF CAREER award 1553284 and NSF III 1900933. MZ was supported by ODNI, IARPA, via the BETTER Program contract 2019-19051600005. The research of KK was partially supported by the Center of Mathematical Sciences and Applications at Harvard University. The views, opinions, and/or findings contained in this article are those of the author and should not be interpreted as representing the official views or policies, either expressed or implied, of the Defense Advanced Research Projects Agency, the Department of Defense, ODNI, IARPA, or the U.S. Government. The U.S. Government is authorized to reproduce and distribute reprints for governmental purposes notwithstanding any copyright annotation therein.

References

- Abu-El-Haija, S., Kapoor, A., Perozzi, B., and Lee, J. N-gcn: Multi-scale graph convolution for semi-supervised node classification. In *Uncertainty in Artificial Intelligence*, pp. 841–851. PMLR, 2020.
- Allen-Zhu, Z., Li, Y., and Song, Z. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pp. 242–252, 2019.
- Arora, S., Cohen, N., and Hazan, E. On the optimization of deep networks: Implicit acceleration by overparameterization. In *International Conference on Machine Learning*, pp. 244–253. PMLR, 2018.
- Arora, S., Cohen, N., Golowich, N., and Hu, W. A convergence analysis of gradient descent for deep linear neural networks. In *International Conference on Learning Representations*, 2019a.
- Arora, S., Du, S., Hu, W., Li, Z., and Wang, R. Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks. In *International Conference on Machine Learning*, pp. 322–332, 2019b.
- Bartlett, P. L., Helmbold, D. P., and Long, P. M. Gradient descent with identity initialization efficiently learns positive-definite linear transformations by deep residual networks. *Neural computation*, 31(3):477–502, 2019.
- Cai, T., Luo, S., Xu, K., He, D., Liu, T.-y., and Wang, L. Graphnorm: A principled approach to accelerating graph neural network training. *arXiv preprint arXiv:2009.03294*, 2020.
- Chen, J., Ma, T., and Xiao, C. FastGCN: Fast learning with graph convolutional networks via importance sampling. In *International Conference on Learning Representations*, 2018a.
- Chen, J., Zhu, J., and Song, L. Stochastic training of graph convolutional networks with variance reduction. In *International Conference on Machine Learning*, pp. 942–950, 2018b.
- Chen, M., Wei, Z., Huang, Z., Ding, B., and Li, Y. Simple and deep graph convolutional networks. In *International Conference on Machine Learning*, pp. 1725–1735. PMLR, 2020.
- Chen, Z., Villar, S., Chen, L., and Bruna, J. On the equivalence between graph isomorphism testing and function approximation with gnns. In *Advances in Neural Information Processing Systems*, pp. 15894–15902, 2019.
- Chiang, W.-L., Liu, X., Si, S., Li, Y., Bengio, S., and Hsieh, C.-J. Cluster-gcn: An efficient algorithm for training deep and large graph convolutional networks. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 257–266, 2019.
- Chizat, L., Oyallon, E., and Bach, F. On lazy training in differentiable programming. In *Advances in Neural Information Processing Systems*, pp. 2937–2947, 2019.
- Defferrard, M., Bresson, X., and Vandergheynst, P. Convolutional neural networks on graphs with fast localized spectral filtering. In *Advances in Neural Information Processing Systems*, pp. 3844–3852, 2016.
- Du, S. and Hu, W. Width provably matters in optimization for deep linear neural networks. In *International Conference on Machine Learning*, pp. 1655–1664, 2019.
- Du, S., Lee, J., Li, H., Wang, L., and Zhai, X. Gradient descent finds global minima of deep neural networks. In *International Conference on Machine Learning*, pp. 1675–1685, 2019a.
- Du, S. S., Hou, K., Salakhutdinov, R. R., Póczos, B., Wang, R., and Xu, K. Graph neural tangent kernel: Fusing graph neural networks with graph kernels. In *Advances in Neural Information Processing Systems*, pp. 5724–5734, 2019b.
- Du, S. S., Zhai, X., Póczos, B., and Singh, A. Gradient descent provably optimizes over-parameterized neural networks. In *International Conference on Learning Representations*, 2019c.
- Duvenaud, D. K., Maclaurin, D., Iparraguirre, J., Bombarell, R., Hirzel, T., Aspuru-Guzik, A., and Adams, R. P. Convolutional networks on graphs for learning molecular fingerprints. In *Advances in neural information processing systems*, pp. 2224–2232, 2015.

- Fey, M. and Lenssen, J. E. Fast graph representation learning with pytorch geometric. *arXiv preprint arXiv:1903.02428*, 2019.
- Garg, V. K., Jegelka, S., and Jaakkola, T. Generalization and representational limits of graph neural networks. In *International Conference on Machine Learning*, 2020.
- Gilmer, J., Schoenholz, S. S., Riley, P. F., Vinyals, O., and Dahl, G. E. Neural message passing for quantum chemistry. In *International Conference on Machine Learning*, pp. 1273–1272, 2017.
- Gori, M., Monfardini, G., and Scarselli, F. A new model for learning in graph domains. In *Proceedings. 2005 IEEE International Joint Conference on Neural Networks, 2005.*, volume 2, pp. 729–734. IEEE, 2005.
- Hamilton, W. L., Ying, R., and Leskovec, J. Inductive representation learning on large graphs. In *Advances in Neural Information Processing Systems*, pp. 1025–1035, 2017.
- Hardt, M. and Ma, T. Identity matters in deep learning. In *International Conference on Learning Representations*, 2017.
- Huang, J. and Yau, H.-T. Dynamics of deep neural networks and neural tangent hierarchy. In *International conference on machine learning*, 2020.
- Huang, W., Zhang, T., Rong, Y., and Huang, J. Adaptive sampling towards fast graph representation learning. *Advances in neural information processing systems*, 31: 4558–4567, 2018.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems*, pp. 8571–8580, 2018.
- Ji, Z. and Telgarsky, M. Directional convergence and alignment in deep learning. *arXiv preprint arXiv:2006.06657*, 2020.
- Kawaguchi, K. Deep learning without poor local minima. In *Advances in Neural Information Processing Systems*, pp. 586–594, 2016.
- Kawaguchi, K. On the theory of implicit deep learning: Global convergence with implicit layers. In *International Conference on Learning Representations (ICLR)*, 2021.
- Kawaguchi, K. and Huang, J. Gradient descent finds global minima for generalizable deep neural networks of practical sizes. In *2019 57th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, pp. 92–99. IEEE, 2019.
- Kearnes, S., McCloskey, K., Berndl, M., Pande, V., and Riley, P. Molecular graph convolutions: moving beyond fingerprints. *Journal of computer-aided molecular design*, 30(8):595–608, 2016.
- Keriven, N. and Peyré, G. Universal invariant and equivariant graph neural networks. In *Advances in Neural Information Processing Systems*, pp. 7092–7101, 2019.
- Kipf, T. N. and Welling, M. Semi-supervised classification with graph convolutional networks. In *International Conference on Learning Representations*, 2017.
- Laurent, T. and Brecht, J. Deep linear networks with arbitrary loss: All local minima are global. In *International conference on machine learning*, pp. 2902–2907. PMLR, 2018.
- Lee, J., Xiao, L., Schoenholz, S., Bahri, Y., Novak, R., Sohl-Dickstein, J., and Pennington, J. Wide neural networks of any depth evolve as linear models under gradient descent. In *Advances in neural information processing systems*, pp. 8572–8583, 2019.
- Li, G., Muller, M., Thabet, A., and Ghanem, B. Deepgcns: Can gcns go as deep as cnns? In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 9267–9276, 2019.
- Li, G., Xiong, C., Thabet, A., and Ghanem, B. Deep-ergcn: All you need to train deeper gcns. *arXiv preprint arXiv:2006.07739*, 2020.
- Li, Y. and Liang, Y. Learning overparameterized neural networks via stochastic gradient descent on structured data. In *Advances in Neural Information Processing Systems*, pp. 8157–8166, 2018.
- Liao, P., Zhao, H., Xu, K., Jaakkola, T., Gordon, G., Jegelka, S., and Salakhutdinov, R. Graph adversarial networks: Protecting information against adversarial attacks. *arXiv preprint arXiv:2009.13504*, 2020.
- Liu, C., Zhu, L., and Belkin, M. On the linearity of large non-linear models: when and why the tangent kernel is constant. *Advances in Neural Information Processing Systems*, 33, 2020.
- Loukas, A. How hard is to distinguish graphs with graph neural networks? In *Advances in neural information processing systems*, 2020.
- Maron, H., Ben-Hamu, H., Serviansky, H., and Lipman, Y. Provably powerful graph networks. In *Advances in Neural Information Processing Systems*, pp. 2156–2167, 2019.

- Merkwirth, C. and Lengauer, T. Automatic generation of complementary descriptors with molecular graph networks. *J. Chem. Inf. Model.*, 45(5):1159–1168, 2005.
- Nitanda, A. and Suzuki, T. Optimal rates for averaged stochastic gradient descent under neural tangent kernel regime. In *International Conference on Learning Representations*, 2021.
- Oono, K. and Suzuki, T. Optimization and generalization analysis of transduction through gradient boosting and application to multi-scale graph neural networks. *Advances in Neural Information Processing Systems*, 33, 2020.
- Sato, R., Yamada, M., and Kashima, H. Approximation ratios of graph neural networks for combinatorial problems. In *Advances in Neural Information Processing Systems*, pp. 4083–4092, 2019.
- Sato, R., Yamada, M., and Kashima, H. Random features strengthen graph neural networks. *arXiv preprint arXiv:2002.03155*, 2020.
- Saxe, A. M., McClelland, J. L., and Ganguli, S. Exact solutions to the nonlinear dynamics of learning in deep linear neural networks. In *International Conference on Learning Representations*, 2014.
- Scarselli, F., Gori, M., Tsoi, A. C., Hagenbuchner, M., and Monfardini, G. The graph neural network model. *IEEE Transactions on Neural Networks*, 20(1):61–80, 2009.
- Scarselli, F., Tsoi, A. C., and Hagenbuchner, M. The vapnik–chervonenkis dimension of graph and recursive neural networks. *Neural Networks*, 108:248–259, 2018.
- Sen, P., Namata, G., Bilgic, M., Getoor, L., Galligher, B., and Eliassi-Rad, T. Collective classification in network data. *AI magazine*, 29(3):93, 2008.
- Thekumparampil, K. K., Wang, C., Oh, S., and Li, L.-J. Attention-based graph neural network for semi-supervised learning. *arXiv preprint arXiv:1803.03735*, 2018.
- Velickovic, P., Cucurull, G., Casanova, A., Romero, A., Lio, P., and Bengio, Y. Graph attention networks. In *International Conference on Learning Representations*, 2018.
- Vignac, C., Loukas, A., and Frossard, P. Building powerful and equivariant graph neural networks with message-passing. *Advances in neural information processing systems*, 2020.
- Wu, F., Souza, A., Fifty, C., Yu, T., and Weinberger, K. Simplifying graph convolutional networks. In *International Conference on Machine Learning*, 2019.
- Xu, K., Li, C., Tian, Y., Sonobe, T., Kawarabayashi, K.-i., and Jegelka, S. Representation learning on graphs with jumping knowledge networks. In *International Conference on Machine Learning*, pp. 5453–5462, 2018.
- Xu, K., Hu, W., Leskovec, J., and Jegelka, S. How powerful are graph neural networks? In *International Conference on Learning Representations*, 2019.
- Xu, K., Li, J., Zhang, M., Du, S. S., ichi Kawarabayashi, K., and Jegelka, S. What can neural networks reason about? In *International Conference on Learning Representations*, 2020.
- Xu, K., Zhang, M., Li, J., Du, S. S., Kawarabayashi, K.-I., and Jegelka, S. How neural networks extrapolate: From feedforward to graph neural networks. In *International Conference on Learning Representations*, 2021.
- Yehudai, G. and Shamir, O. On the power and limitations of random features for understanding neural networks. In *Advances in Neural Information Processing Systems*, pp. 6598–6608, 2019.
- Zhang, S., Wang, M., Liu, S., Chen, P.-Y., and Xiong, J. Fast learning of graph neural networks with guaranteed generalizability: One-hidden-layer case. In *International Conference on Machine Learning*, pp. 11268–11277, 2020.
- Zou, D., Hu, Z., Wang, Y., Jiang, S., Sun, Y., and Gu, Q. Layer-dependent importance sampling for training deep and large graph convolutional networks. In *Advances in Neural Information Processing Systems*, pp. 11249–11259, 2019.
- Zou, D., Long, P. M., and Gu, Q. On the global convergence of training deep linear resnets. In *International Conference on Learning Representations*, 2020.