



Exploring Over-smoothing in Graph Attention Networks from the Markov Chain Perspective

Weichen Zhao
Chenguang Wang
Congying Han*
Tiande Guo

zhaoweichen14@mails.ucas.ac.cn
wangchenguang19@mails.ucas.ac.cn
hancy@ucas.ac.cn
tdguo@ucas.ac.cn

School of Mathematical Sciences University of Chinese Academy of Sciences (UCAS)
Shijingshan Qu, Beijing Shi, China

ABSTRACT

The over-smoothing problem causing the depth limitation is an obstacle of developing deep graph neural network (GNN). Compared with Graph Convolutional Networks (GCN), over-smoothing in Graph Attention Network (GAT) has not drawn enough attention. In this work, we analyze the over-smoothing problem in GAT from the Markov chain perspective. First we establish a connection between GAT and a time-inhomogeneous random walk on the graph. Then we show that the GAT is not always over-smoothing using conclusions in the time-inhomogeneous Markov chain. Finally, we derive a sufficient condition for GAT to avoid over-smoothing in the Markovian sense based on our findings about the existence of the limiting distribution of the time-inhomogeneous Markov chain. We design experiments to verify our theoretical findings. Results show that our proposed sufficient condition can effectively improve over-smoothing problem in GAT and enhance the performance of the model.

CCS CONCEPTS

• **Computing methodologies** → **Neural networks**; *Supervised learning by classification*; *Regularization*.

KEYWORDS

Graph neural networks, Graph attention networks, Over-smoothing, Markov chains

ACM Reference Format:

Weichen Zhao, Chenguang Wang, Congying Han, and Tiande Guo. 2023. Exploring Over-smoothing in Graph Attention Networks from the Markov Chain Perspective. In *2023 International Conference on Frontiers of Artificial Intelligence and Machine Learning (FAIML 2023)*, April 14–16, 2023, Beijing.

*Corresponding author at: No.19A, Yuquan Road, Shijingshan District, Beijing.



This work is licensed under a Creative Commons Attribution International 4.0 License.

FAIML 2023, April 14–16, 2023, Beijing, China
© 2023 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-0754-4/23/04.
<https://doi.org/10.1145/3616901.3617022>

China. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3616901.3617022>

1 INTRODUCTION

Graph neural networks [3, 7, 12, 20] have achieved great success in processing graph data which is rich in information about the relationships between objects. GAT [20] is one of the most representative GNN models. It introduces the attention mechanism into GNN and inspires a class of attention-based GNN models [1, 13, 23].

The deepening of the network has brought about changes in neural networks and caused a boom in deep learning. Unlike typical deep neural networks, in the training of graph neural networks, researchers have found that the performance of GNN decreases instead as the depth increases. There are several possible reasons for the depth limitations of GNN. Li et al. [16] first attribute this anomaly to *over-smoothing*, a phenomenon in which the representations of different nodes tend to be consistent as the network deepens, leading to indistinguishable node representations. Many researchers have studied this problem and proposed some improvement methods [4–6, 16–18, 24]. However, the research of the over-smoothing problem has mostly focused on graph convolutional network. There is a lack of unique analysis of over-smoothing in GAT and corresponding improvement methods.

Noting the Markov property of the forward propagation process of GNNs and considering the node set as a state space, in this work, we connect GAT with a time-inhomogeneous random walk on the graph. Considering the nodes' representations as distributions on the state space, we interpret the over-smoothing in GAT as the convergence of the representation distribution to the limiting distribution. Using conclusions of the time-inhomogeneous Markov chain, we show that GAT does not necessarily suffer from over-smoothing in the Markovian sense. Further, we prove a necessary condition for the existence of the limiting distribution. Based on this conclusion, we derive a sufficient condition for GAT to avoid over-smoothing.

We verify our conclusions on the benchmark datasets. Based on the sufficient condition, we propose a regularization term which can be flexibly added to the training of the neural network. Results show that our proposed sufficient condition can significantly improve the performance of GAT. In addition, the representation learned by

different nodes is more inconsistent after adding the regularization term, which indicates that the over-smoothing in GAT is improved.

Contributions. In summary, our contributions are as follows:

- We establish a connection between GAT and a time-inhomogeneous random walk on the graph. Then we show that GAT is not always over-smoothing in the Markovian sense (Section 3).
- We study the existence of limiting distributions of the time-inhomogeneous Markov chain. And based on this, we give a sufficient condition for GAT to avoid over-smoothing (Section 4 and Theorem 4.1, 4.2).
- We propose a regularization term based on this sufficient condition, which can be simply and flexibly added to the training of GAT, and experimentally verify that our proposed condition can improve the model performance by solving the over-smoothing problem of GAT (Section 5).

Notation. Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be a connected non-bipartite graph, where $\mathcal{V} := \{1, 2, \dots, N\}$ is the node set, $\mathcal{E}$ is the edge set, $N = |\mathcal{V}|$ is the number of nodes. If there are connected edges between nodes $u, v \in \mathcal{V}$, then denote by $(u, v) \in \mathcal{E}$. $\deg(u)$ denotes the degree of node $u \in \mathcal{V}$ and $\mathcal{N}(v)$ denotes the neighbors of node v . The corresponding adjacency matrix is A and the degree matrix is D . Let $(\Omega, \mathcal{F}, \mathbf{P})$ be a probability space, $(\mathcal{M}, \mathcal{M})$ be a finite state space. $\vec{x} = \{x_n, n \in \mathcal{T}\}$ is a stochastic process taking values in $\mathcal{M}$. $\mathcal{T}$ is a time parameter set. $P(i, j)$ denotes the element i, j of matrix P .

2 RELATED WORK

In this section, we introduce graph attention network and the related work of over-smoothing.

2.1 Graph Attention Network

GAT [20] establishes attention functions between nodes u and v with connected edges $(u, v) \in \mathcal{E}$

$$\alpha_{u,v}^{(l)} = \frac{\exp(\phi^{(l)}(h_u^{(l-1)}, h_v^{(l-1)}))}{\sum_{k \in \mathcal{N}(u)} \exp(\phi^{(l)}(h_u^{(l-1)}, h_k^{(l-1)}))} \quad (1)$$

where $h_u^{(l)} \in \mathbb{R}^F$ is the embedding for node u at the layer l and $\phi^{(l)}(h_u^{(l-1)}, h_v^{(l-1)}) := \text{LeakyReLU}(\mathbf{a}^T [W^{(l)} h_u^{(l-1)} \parallel W^{(l)} h_v^{(l-1)}])$, where $\mathbf{a} \in \mathbb{R}^{2F}$ and $W^{(l)}$ is the weight matrix. Then the GAT layer is defined as

$$h_u^{(l)} := \sigma_{W^{(l)}} \left(\sum_{v \in \mathcal{N}(u)} \alpha_{u,v}^{(l)} h_v^{(l-1)} \right).$$

Written in matrix form

$$H^{(l)} = \sigma_{W^{(l)}} (P_{\text{att}}^{(l)} H^{(l-1)}),$$

where σ is the activation function, $P_{\text{att}}^{(l)} \in \mathbb{R}^{N \times N}$ is the attention matrix satisfying $P_{\text{att}}^{(l)}(u, v) = \alpha_{u,v}^{(l)}$ if $v \in \mathcal{N}(u)$, otherwise $P_{\text{att}}^{(l)}(u, v) = 0$ and $\sum_{v=1}^N P_{\text{att}}^{(l)}(u, v) = 1$.

2.2 Over-smoothing

There is a phenomenon that the GNN has better experimental results in the shallow layer case, and instead do not work well in the deep layer case. The researchers find that this is due to the fact that during the GNN training process, the hidden layer representation of each node tend to converge to the same value as the number of

layers increases. This phenomenon is called over-smoothing. This problem affects the deepening of GNN layers and limits the further development of GNN.

Intuitively, Zhao & Akoglu [24] proposes a normalization layer, Pairnorm, to avoid node representations from becoming too similar. Thus, the over-smoothing phenomenon is alleviated.

Another intuitive analysis of the over-smoothing is that as the network is stacked, the model forgets the initial input features and only updates the representations based on the structure of the graph data. It is natural to think that the problem of the model forgetting the initial features can be alleviated by reminding the network what its previous features are. Many methods have been proposed based on such intuitive analysis. The simplest one, Kipf & Welling [12] propose to add residual connections to graph convolutional networks. The node representations of the hidden layer l are directly added to the node representations of the previous layer to remind the network not to forget the previous features. However, Chiang et al. [6] argues that residual connectivity ignores the structure of the graph and should be considered to reflect more the influence of the weights of different neighboring nodes. So this work gives more weight to the representations from the previous layer in the message passing of each GCN layer by improving the graph convolution operator. Chen et al. [5]; Li et al. [14]; Xu et al. [22] also use this idea.

Oono & Suzuki [17] connects the GCN with a dynamical system and interprets the over-smoothing problem as the convergence of the dynamical system to an invariant subspace. Rong et al. [18] proposed DropEdge method based on the perspective of dynamical system. The idea of DropEdge is to randomly drop some edges in the original graph at each layer. This operation slows down the convergence of the dynamical system to the invariant subspace. Thus DropEdge method can alleviate the over-smoothing.

Most of the works on over-smoothing focus on GCN and ignore the discussion of GAT. Wang et al. [21] first analyze the over-smoothing problem in GAT and improve GAT via margin-based constraints. However, we disagree with their conclusion that GAT will be over-smoothing. We discuss this in detail in Section 3.

3 ANALYSIS OF OVER-SMOOTHING IN GAT

In this section, we analyze the over-smoothing problem in GAT from the Markov chain perspective. We show that forward propagation of GAT is a time-inhomogeneous random walk $\vec{v}_{\text{att}}$ on the graph, and that over-smoothing is caused by the convergence of the representation distribution to the limiting distribution. Next, we show that GAT is not always oversmooth by analyzing that the limiting distribution of $\vec{v}_{\text{att}}$ does not always exist.

3.1 Relationship between GAT and time-inhomogeneous random walk

We first connect GAT with a time-inhomogeneous random walk on the graph. The following defines the general random walk on the graph.

Definition 3.1. Given a graph $\mathcal{G}$ and a starting node $u \in \mathcal{V}$, we select a neighbor $v \in \mathcal{N}(u)$ of it with positive probability

$P^{(1)}(u, v) > 0$, and move to its neighbor. $P^{(1)}(u, v)$ satisfies

$$\sum_{v \in \mathcal{N}(u)} P^{(1)}(u, v) = 1.$$

Then we repeat this process. $P^{(t)}(u, v)$, $t = 1, 2, \dots$ is not always the same. The random sequence of nodes selected this way is a *random walk on the graph*.

Recalling the definition of GAT, we focus on its message passing $H^{(l)} = P_{\text{att}}^{(l)} H^{(l-1)}$. Since $P_{\text{att}}^{(l)}(u, v) \geq 0$ and $\sum_{v=1}^N P_{\text{att}}^{(l)}(u, v) = 1$, $P_{\text{att}}^{(l)}$ is a one-step stochastic matrix of a random walk on the graph. Moreover, since $P_{\text{att}}^{(l)}$, $l = 1, 2, \dots$ is not the same, the forward propagation process of node representations in GAT is a time-inhomogeneous random walk on a graph, denoted as $\vec{v}_{\text{att}}$. It has the state space $\mathcal{V}$ and the family of stochastic matrices $\{P_{\text{att}}^{(1)}, P_{\text{att}}^{(2)}, \dots, P_{\text{att}}^{(l)}, \dots\}$. The inconsistency of the nodes message passing at each layer in GAT causes the time-inhomogeneity of the corresponding chain $\vec{v}_{\text{att}}$, which is an important difference between GAT and GNNs with consistent message passing such as GCN.

Following is the general definition of the limiting distribution. We use it to explain the over-smoothing problem.

Definition 3.2. Let $\vec{x} = \{x_n, n \in \mathcal{T}\}$ be a time-inhomogeneous Markov chain on a finite state space $\mathcal{M}$, the initial distribution be μ_0 and the distribution of the chain $\vec{x}$ at moment n be μ_n . π is the *limiting distribution* of the chain $\vec{x}$, if $\mu_n \rightarrow \pi$, $n \rightarrow \infty$.

The node representation $h = \{h_u, u \in \mathcal{V}\}$ is viewed as a discrete probability distribution over the node set $\mathcal{V}$. If $\vec{v}_{\text{att}}$ has a limiting distribution π , then as the GAT propagates forward, the representation distribution converges to the limiting distribution. This causes the potential over-smoothing problem in GAT. However, the limiting distribution of the time-inhomogeneous Markov chain does not always exist. We next discuss specifically the possibility of GAT to avoid over-smoothing in Markovian sense.

3.2 GAT is not always over-smoothing

In this subsection, we first show that the conclusion that the GAT will be over-smoothing cannot be proven.

The following theorem gives property of the family of stochastic matrices $\{P_{\text{att}}^{(1)}, P_{\text{att}}^{(2)}, \dots, P_{\text{att}}^{(l)}, \dots\}$, shows the existence of stationary distribution of each graph attention matrix, and gives the explicit expression of stationary distribution.

THEOREM 3.3. *There exists a unique probability distribution $\pi^{(l)}$ on $\mathcal{V}$ satisfies $\pi^{(l)} = \pi^{(l)} P_{\text{att}}^{(l)}$, $l = 1, 2, \dots$, where $\pi^{(l)}(u) = \frac{\deg^{(l)}(u)}{\sum_{k \in \mathcal{V}} \deg^{(l)}(k)}$, $\deg^{(l)}(u) = \sum_{z \in \mathcal{N}(u)} \exp(\phi^{(l)}(h_u^{(l-1)}, h_z^{(l-1)}))$.*

Previously, Wang et al. [21] discussed the over-smoothing problem in GAT and concluded that the GAT would over-smooth. Same as our work, they viewed the $P_{\text{att}}^{(l)}$ at each layer as stochastic matrix of a random walk on the graph. However, they ignore the fact that the complete forward propagation process of GAT is essentially a time-inhomogeneous random walk on the graph. The core theorem

¹Similar to Li et al. [16]; Wang et al. [21], we omit the activation function.

stating that the GAT will over-smooth in their work is flawed. In its proof, the stationary distribution $\pi^{(l)}$ of the graph attention matrix $P_{\text{att}}^{(l)}$ for each layer is consistent, i.e.

$$\pi^{(1)} = \pi^{(2)} = \dots = \pi^{(l)} = \dots$$

However, since each layer $\phi^{(l)}$ is different, by Theorem 3.3,

$$\pi^{(1)} \neq \pi^{(2)} \neq \dots \neq \pi^{(l)} \neq \dots$$

The conclusion that the GAT will be over-smoothing cannot be proven. Then we show that the GAT is not always over-smoothing.

Since over-smoothing in GAT is related to the limiting distribution of time-inhomogeneous random walk on the graph, we next focus on the limiting distribution of $\vec{v}_{\text{att}}$.

Compared to the time-homogeneous Markov chain, it is much more difficult to investigate the limiting distribution of the time-inhomogeneous chain. In order to study the convergence of the probability distribution on the state space, we introduce the Dobrushin contraction coefficient and the Dobrushin inequality (Lemma 3.4). See [8] for proof.

LEMMA 3.4. *Let μ and ν be probability distributions on a finite state space $\mathcal{M}$ and P be a stochastic matrix, then*

$$\|\mu P - \nu P\|_1 \leq C(P) \|\mu - \nu\|_1,$$

where $C(P) := \frac{1}{2} \sup_{i,j} \sum_{k \in \mathcal{M}} |P(i, k) - P(j, k)|$ is called the Dobrushin contraction coefficient of the stochastic matrix P .

Bowerman et al. [2]; Huang et al. [10] discussed the limiting case that an arbitrary initial distribution transfer according to a time-inhomogeneous chain. The sufficient condition for the existence of the limiting distribution is summarized in the following lemma. See [8] for proof.

LEMMA 3.5. *Let $\vec{x} = \{x_n, n \in \mathcal{T}\}$ be a time-inhomogeneous Markov chain on a finite state space $\mathcal{M}$ with stochastic matrix $P^{(n)}$. If the following (1), (2) and (3A) or (3B) are satisfied*

- (1) *There exists a stationary distribution $\pi^{(n)}$ when $P^{(n)}$ is treated as a stochastic matrix of a time-homogeneous chain;*
- (2) $\sum_n \|\pi^{(n)} - \pi^{(n+1)}\|_1 < \infty$;
- (3A) *(Isaacson-Madsen condition) For any probability distribution μ, ν on $\mathcal{M}$ and positive integer k*

$$\|(\mu - \nu)P^{(k)} \dots P^{(n)}\|_1 \rightarrow 0, \quad n \rightarrow \infty.$$

- (3B) *(Dobrushin condition) For any positive integers k*

$$C(P^{(k)} \dots P^{(n)}) \rightarrow 0, \quad n \rightarrow \infty.$$

Then there exists a probability measure π on $\mathcal{M}$ such that

- (1) $\|\pi^{(n)} - \pi\|_1 \rightarrow 0, \quad n \rightarrow \infty$;
- (2) *Let the initial distribution be μ_0 and the distribution of the chain $\vec{x}$ at step n be $\mu_n := \mu_{n-1}P^{(n)}$, then for any initial distribution μ_0 , we have*

$$\|\mu_n - \pi\|_1 \rightarrow 0, \quad n \rightarrow \infty,$$

Returning to GAT, for time-inhomogeneous random walk $\vec{v}_{\text{att}}$, its family of stochastic matrices $\{P_{\text{att}}^{(l)}\}$ satisfies the condition (1) (Theorem 3.3). However, the series of positive terms $\sum_l \|\pi^{(l)} - \pi^{(l+1)}\|_1$ is possible to be divergent and the condition (2) of Lemma 3.5 can not be guaranteed. Moreover, according to the definition of $\{P_{\text{att}}^{(l)}\}$

Table 1: Results of GAT

datasets	model	#layers					
		3	4	5	6	7	8
Cora	GAT	0.7773(± 0.0054)	0.7602(± 0.0166)	0.4821(± 0.3021)	0.2774(± 0.2542)	0.1672(± 0.0780)	0.0958(± 0.0059)
	GAT-RT	0.7884(± 0.0157)	0.7872(± 0.0127)	0.7648(± 0.0077)	0.6454(± 0.2508)	0.3244(± 0.2465)	0.1678(± 0.0756)
Citeseer	GAT	0.6643(± 0.0063)	0.6541(± 0.0076)	0.3472(± 0.2582)	0.2474(± 0.1947)	0.1768(± 0.0064)	0.1902(± 0.0598)
	GAT-RT	0.6678(± 0.0157)	0.6692(± 0.0072)	0.6208(± 0.0380)	0.2706(± 0.1884)	0.1915(± 0.0200)	0.1864(± 0.0229)
Pubmed	GAT	0.7616(± 0.0115)	0.7534(± 0.0114)	0.7653(± 0.0072)	0.7468(± 0.0084)	0.7468(± 0.0045)	0.7076(± 0.0112)
	GAT-RT	0.7673(± 0.0064)	0.7659(± 0.0123)	0.7684(± 0.0063)	0.7664(± 0.0092)	0.7596(± 0.0107)	0.7618(± 0.0114)
ogbn-arxiv	GAT	0.7117(± 0.0023)	0.7144(± 0.0015)	0.7061(± 0.0082)	0.6396(± 0.1031)	0.4307(± 0.1720)	-
	GAT-RT	0.7115(± 0.0025)	0.7104(± 0.0022)	0.7063(± 0.0040)	0.6709(± 0.0344)	0.5653(± 0.1122)	-

(Equation 1), neither the Isaacson-Madsen condition nor the Dobrushin condition can be guaranteed. So the time-inhomogeneous chain $\vec{v}_{\text{att}}$ does not always have a limiting distribution. This indicates that GAT is not always over-smoothing.

4 SUFFICIENT CONDITION FOR GAT TO AVOID OVER-SMOOTHING

In this section, we propose and prove a necessary condition for the existence of limiting distribution for a time-inhomogeneous Markov chain. Then we apply this theoretical result to GAT and propose a sufficient condition to ensure that GAT can avoid over-smoothing.

In the study of Markov chains, researchers usually focus on the sufficient conditions for the existence of the limiting distribution. And the case when the limiting distribution does not exist has rarely been studied. We study the necessary conditions for the existence of the limit distribution in order to obtain sufficient conditions for its nonexistence.

The following theorem gives a necessary condition for the existence of the limit distribution of the time-inhomogeneous Markov chain. Although other necessary conditions exist, Theorem 4.1 is one of the most intuitive and simplest in form.

THEOREM 4.1. *Let $\vec{x} = \{x_n, n \in \mathcal{T}\}$ be a time-inhomogeneous Markov chain on a finite state space $\mathcal{M}$, and write its n -step stochastic matrix as $P^{(n)}$, satisfying that, $P^{(n)}$ is irreducible and aperiodic, there exists a unique stationary distribution $\pi^{(n)}$ as the time-homogeneous transition matrix, and $C(P^{(n)}) < 1$. Let the initial distribution be μ_0 and the distribution of the chain $\vec{x}$ at step n be $\mu_n := \mu_{n-1}P^{(n)}$. Then*

$$\|\pi^{(n)} - \pi\| \rightarrow 0, \quad n \rightarrow \infty$$

is a necessary condition for existence of a probability distribution π on $\mathcal{M}$ such that $\|\mu_n - \pi\| \rightarrow 0, n \rightarrow \infty$.

We explain Theorem 4.1 intuitively. In the limit sense, transition of μ_{n-1} satisfies $\lim_{n \rightarrow \infty} \mu_{n-1}P^{(n)} = \lim_{n \rightarrow \infty} \mu_n = \lim_{n \rightarrow \infty} \mu_{n-1} = \pi$. On the other hand, for all $n > 0$, $\pi^{(n)}$ is the unique solution of the equation $\mu = \mu P^{(n)}$. Thus $\lim_{n \rightarrow \infty} \pi^{(n)} = \lim_{n \rightarrow \infty} \mu_{n-1} = \pi$.

By Theorem 4.1, we give the following sufficient condition for GAT to avoid over-smoothing in Markovian sense.

THEOREM 4.2. *Let $h_u^{(l)}$ be representation of node $u \in \mathcal{V}$ at the hidden layer l in GAT. The sufficient condition for GAT to avoid over-smoothing is that there exists $\delta > 0$ such that for any $l \geq 2$, satisfying*

$$\|h_u^{(l-1)} - h_u^{(l)}\|_1 > \delta. \quad (2)$$

When Equation 2 is satisfied, the time-inhomogeneous random walk $\vec{x}_{\text{att}}$ corresponding to GAT does not have a limiting distribution, and thus GAT avoids potential over-smoothing problems in a Markovian sense. Theorem 4.2 has an intuitive meaning. The essence of over-smoothing is that the node representations converge with the propagation of the network. By Cauchy's convergence test, the condition exactly avoid representation $h_u^{(l)}$ of the node u from converging as network deepens.

Since Theorem 4.1 generally holds for all time-inhomogeneous Markov chains, GNN related to a time-inhomogeneous Markov chain such as GEN [15] can obtain the sufficient conditions similar to Theorem 4.2 to avoid over-smoothing in Markovian sense.

Considering that this sufficient condition is task-agnostic, we can formulate this condition to a regularization term and add it to the loss function determined by its original task. Formally, assume the original loss function is $L_\theta(x)$, the total loss function can be rewritten as:

$$\hat{L}_\theta(x) = L_\theta(x) + \text{RT}_\theta(x) \quad (3)$$

where θ is the parameter of neural networks and $\text{RT}_\theta(x)$ is the regularization term determined by the sufficient condition in Theorem 4.2.

5 EXPERIMENTS

In this section, we experimentally verify the correctness of our theoretical results. We rewrite the sufficient condition for GAT to avoid over-smoothing in the Theorem 4.2 as a regularization term. It can be flexibly added to the training of the network. The experimental results show that our proposed condition can effectively avoid the over-smoothing problem and improve the performance of GAT.

5.1 Setup

In this section we briefly introduce the experimental settings. See Appendix A for more specific settings. We verify our conclusions while keeping the other hyperparameters the same² (network structure, learning rate, dropout, epoch, etc.).

²Since we do not aim to refresh State of the Arts, these are not necessarily the optimal hyperparameters.

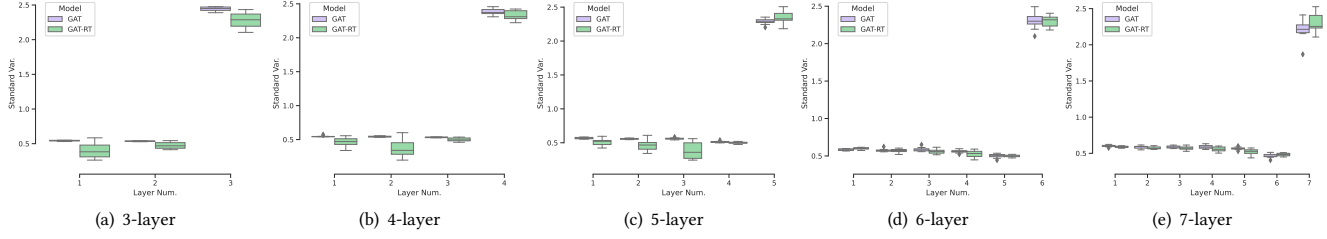


Figure 1: Measurement of over-smoothing of GAT on ogbn-arxiv.

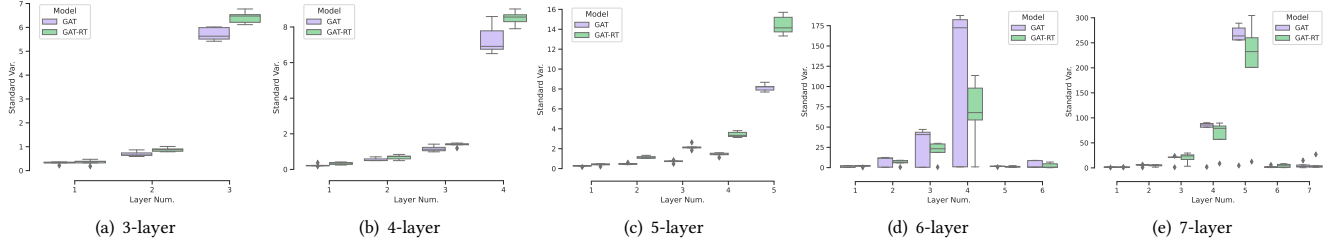


Figure 2: Measurement of over-smoothing of GAT on Cora.

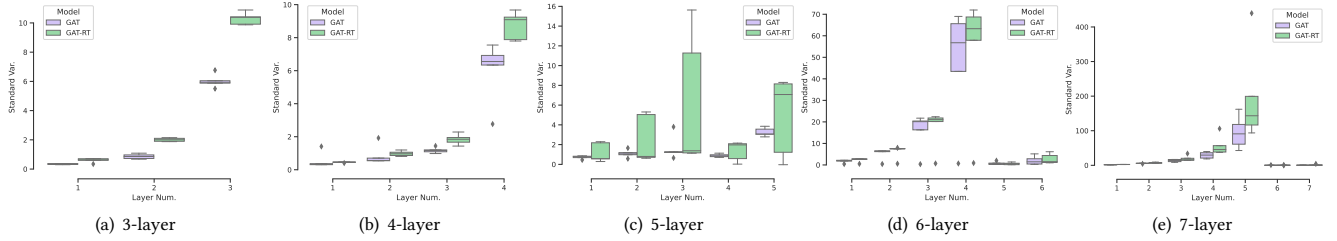


Figure 3: Measurement of over-smoothing of GAT on Citeseer.

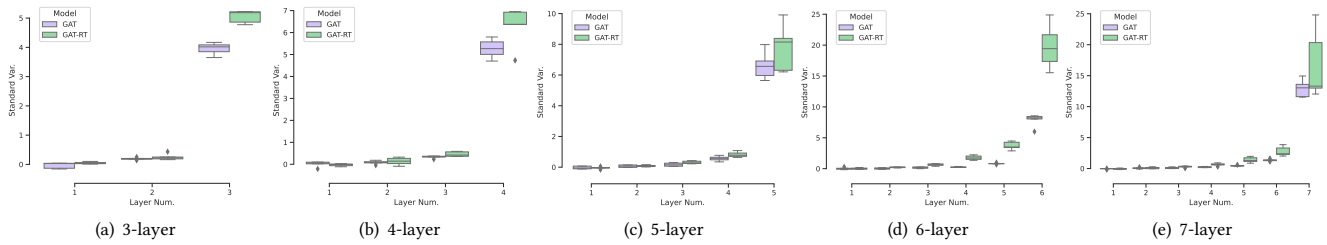


Figure 4: Measurement of over-smoothing of GAT on Pubmed.

Variant of sufficient condition. Notice that the sufficient condition in the Theorem 4.2 is in the form of inequality, which is not conducive to experiments. In the concrete implementation, let $h_u^{(l)}$ be representation of node $u \in \mathcal{V}$ at the hidden layer l . We normalize the distance of the node representations between two adjacent layers and then let it approximate to a given hyperparameter threshold $T \in (0, 1)$, i.e., for the GNN model with n layers, we

obtain a regularization term:

$$\text{RT}_{\theta}(x) = \left(\frac{1}{n} \sum_{l=1}^n (\| \text{Sigmoid}(h_u^{(l-1)}) - \text{Sigmoid}(h_u^{(l)}) \|) - T \right)^2. \quad (4)$$

Since there must exist $\delta > 0$ that satisfies $T > \delta$, Equation 2 can be satisfied if this term is perfectly minimized. For a detailed choice of the threshold T , we put it in Appendix A.

Datasets. In terms of datasets, we follow the datasets used in the original work of GAT [20] as well as the OGB benchmark. We use four standard benchmark datasets - ogbn-arxiv [9], Cora, Citeseer, and Pubmed [19], covering the basic transductive learning tasks.

Implementation details. For the specific implementation, we refer to the open-source code of vanilla GAT, and models with different layers share the same settings: We use the Adam SGD optimizer [11] with learning rate 0.01, the hidden dimension is 64, each GAT layer has 8 heads and the amount of training epoch is 500. All experiments are conducted on a single Nvidia Tesla v100.

5.2 Results of GAT

For record simplicity, we denote the GAT after adding the regularization term to the training as GAT-RT. To keep statistical confidence, we repeat all experiments 10 times and record the mean value and standard deviation. Results shown in Table 1 demonstrate that almost on each dataset and number of layers, GAT-RT will obtain an improvement in the performance. Specifically, on Cora and Citeseer, GAT’s performance begins to decrease drastically when layer numbers surpass 6 and 5 but GAT-RT can relieve this trend in some way. On Pubmed, vanilla GAT’s performance has a gradual decline. The performance of vanilla GAT decrease 6% when the layer number is 8. However, GAT-RT’s performance keeps competitive for all layer numbers. For ogbn-arxiv, GAT-RT performs as competitive as GAT when the layer number is small but outperforms GAT by a big margin when the layer number is large. Specifically when the layer number is 6 and 7, the performance improves roughly by 3% and 13% respectively. Although the performance of GAT-RT decreases with the increase of the number of layers, this is because the performance of the model is also affected by factors such as over-fitting.

5.3 Verification of avoiding over-smoothing

In this subsection, we further experimentally show that the sufficient conditions in Section 4 not only improve the performance of the model but also do avoid the over-smoothing problem.

Since the neural network is a black-box model, we cannot explicitly compute the stationary distribution of the graph neural network when it is over-smoothed. Therefore we measure the degree of over-smoothing by calculating the standard deviation of each node’s representation at each layer. A lower value implies more severe over-smoothing.

Results shown in Fig. 1-4 demonstrate that the node representations obtained from GAT-RT are more diverse than those from GAT, which means the alleviation of over-smoothing. It’s also interesting that there is an accordance between the performance and over-smoothing, for example on Cora dataset, the performance would have a huge decrease when the number of layers is larger than 5, Fig. 2 shows the over-smoothing phenomenon is severe at the same time. Also on Pubmed dataset, the performance is relatively stable and the corresponding Fig. 4 shows that the model trained on this dataset suffers from over-smoothing lightly. These results enlighten us that over-smoothing may be caused by various

objective reasons, e.g. the property of the dataset, and GAT-RT can relieve this negative effect to some extent.

6 CONCLUSION

In this work, we analyze the over-smoothing problem in GAT from a Markov chain perspective. First we relate GAT to a time-inhomogeneous random walk on the graph. By analyzing the limiting distribution of this random walk, we show that it is possible for GAT to avoid potential over-smoothing. Then, we study the limiting distribution of the general time-inhomogeneous Markov chain, and propose a necessary condition for the existence of the limiting distribution. Based on this result, we derive a sufficient condition for GAT to avoid over-smoothing in the Markovian sense. Our results can also be generalized to other GNN models related to time-inhomogeneous Markov chains. Finally, in our experiments we design a regularization term which can be flexibly added to the training. Results on the benchmark datasets show that our theoretical analysis is correct.

The discussion in this paper shows that we can use theoretical results in Markov chains to study problems in GNNs. This motivates us to discuss the over-smoothing problem in more general GNNs through a Markov chain perspective in future work. Further, we can model general GNNs with Markov chains to study more problems other than over-smoothing.

ACKNOWLEDGMENTS

This paper is supported by National key research and development program of China (2021YFA1000403). And it is supported by the National Natural Science Foundation of China (Nos.11991022, U19B2040), and the Strategic Priority Research Program of Chinese Academy of Sciences (No.XDA27000000) and the Fundamental Research Funds for the Central Universities.

A EXPERIMENTAL DETAILS

In this appendix, we add more details on the experiments. Table 2 shows the basic information of each dataset used in our experiments. Table 3 demonstrates the configuration of GNN models, actually, we keep the same setting in the corresponding paper, the only difference is we add the extra proposed regularization term in the optimization objective. In Table 4, we show the detailed selection of threshold T in Equation 4.

Table 2: Summary of the statistics and data split of datasets.

Dataset	(Avg.) Nodes	(Avg.) Edges	Features	Class	Train(#!/%)	Val.(#!/%)	Test(#!/%)
Cora	2708(1 graph)	5429	1433	7	140	500	1000
Citeseer	3327(1 graph)	4732	3703	6	120	500	1000
Pubmed	19717(1 graph)	44338	500	3	60	500	1000
ogbn-arxiv	169,343(1 graph)	1,166,243	128	40	0.54	0.18	0.28

Table 3: Training configuration

Model	Dataset	Hidden.	LR.	Dropout	Epoch
GAT	Cora	64	1e-2	0.5	500
	Citeseer	64	1e-2	0.5	500
	Pubmed	64	1e-2	0.5	500
	ogbn-arxiv	256	1e-2	0.5	500

Table 4: Selection of threshold T on different layer numbers

datasets	model	#layers					
		3	4	5	6	7	8
Cora	GAT-RT	1	0.5	0.6	0.8	1	0.8
Citeseer	GAT-RT	0.3	0.3	1	0.4	0.2	0.1
Pubmed	GAT-RT	0.3	0.2	0.2	0.8	0.5	0.4
ogbn-arxiv	GAT-RT	0.5	0.4	0.7	0.6	0.3	0.5

REFERENCES

- [1] Sami Abu-El-Hajja, Bryan Perozzi, Rami Al-Rfou, and Alexander A Alemi. 2018. Watch your step: Learning node embeddings via graph attention. *Advances in neural information processing systems* 31 (2018).
- [2] Bruce Bowerman, HT David, and Dean Isaacson. 1977. The convergence of Cesaro averages for certain nonstationary Markov chains. *Stochastic processes and their applications* 5, 3 (1977), 221–230.
- [3] Joan Bruna, Wojciech Zaremba, Arthur Szlam, and Yann LeCun. 2013. Spectral Networks and Locally Connected Networks on Graphs. *arXiv preprint arXiv:1312.6203* (2013).
- [4] Deli Chen, Yankai Lin, Wei Li, Peng Li, Jie Zhou, and Xu Sun. 2020. Measuring and relieving the over-smoothing problem for graph neural networks from the topological view. In *Proceedings of the AAAI Conference on Artificial Intelligence*. 3438–3445.
- [5] Ming Chen, Zhewei Wei, Zengfeng Huang, Bolin Ding, and Yaliang Li. 2020. Simple and deep graph convolutional networks. In *International Conference on Machine Learning*. 1725–1735.
- [6] Wei-Lin Chiang, Xuanqing Liu, Si Si, Yang Li, Samy Bengio, and Cho-Jui Hsieh. 2019. 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*. 257–266.
- [7] Michaël Defferrard, Xavier Bresson, and Pierre Vandergheynst. 2016. Convolutional neural networks on graphs with fast localized spectral filtering. *Advances in neural information processing systems* 29 (2016).
- [8] Guanglu Gong and Mingping Qian. 2004. *Tutorial on applied stochastic processes and stochastic models in algorithms and intelligent computing (in Chinese)*. Tsinghua University Press.
- [9] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. 2020. Open Graph Benchmark: Datasets for Machine Learning on Graphs. *arXiv preprint arXiv:2005.00687* (2020).
- [10] Cheng-Chi Huang, Dean Isaacson, and B Vinograd. 1976. The rate of convergence of certain nonhomogeneous Markov chains. *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete* 35, 2 (1976), 141–146.
- [11] Diederik P Kingma and Jimmy Ba. 2015. Adam: A Method for Stochastic Optimization. In *International Conference on Learning Representations*.
- [12] Thomas N Kipf and Max Welling. 2017. Semi-Supervised Classification with Graph Convolutional Networks. In *International Conference on Learning Representations*.
- [13] John Boaz Lee, Ryan Rossi, and Xiangnan Kong. 2018. Graph classification using structural attention. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. 1666–1674.
- [14] Guohao Li, Matthias Müller, Ali Thabet, and Bernard Ghanem. 2019. DeepGCNs: Can GCNs Go As Deep As CNNs?. In *International Conference on Computer Vision*. 9266–9275.
- [15] Guohao Li, Chenxin Xiong, Ali Thabet, and Bernard Ghanem. 2020. Deepergcns: All you need to train deeper gcns. *arXiv preprint arXiv:2006.07739* (2020).
- [16] Qimai Li, Zhichao Han, and Xiao-Ming Wu. 2018. Deeper insights into graph convolutional networks for semi-supervised learning. In *Thirty-Second AAAI conference on artificial intelligence*.
- [17] Kenta Oono and Taiji Suzuki. 2020. Graph Neural Networks Exponentially Lose Expressive Power for Node Classification. In *International Conference on Learning Representations*.
- [18] Yu Rong, Wenbing Huang, Tingyang Xu, and Junzhou Huang. 2020. Droppedge: Towards deep graph convolutional networks on node classification. In *International Conference on Learning Representations*.
- [19] Prithviraj Sen, Galileo Namata, Mustafa Bilgic, Lise Getoor, Brian Galligher, and Tina Eliassi-Rad. 2008. Collective classification in network data. *AI magazine* 29, 3 (2008), 93–93.
- [20] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. 2018. Graph Attention Networks. In *International Conference on Learning Representations*.
- [21] Guangtao Wang, Rex Ying, Jing Huang, and Jure Leskovec. 2019. Improving graph attention networks with large margin-based constraints. *arXiv preprint arXiv:1910.11945* (2019).
- [22] Keyulu Xu, Chengtao Li, Yonglong Tian, Tomohiro Sonobe, Ken-ichi Kawarabayashi, and Stefanie Jegelka. 2018. Representation learning on graphs with jumping knowledge networks. In *International Conference on Machine Learning*. 5453–5462.
- [23] Jiani Zhang, Xingjian Shi, Junyuan Xie, Hao Ma, Irwin King, and Dit Yan Yeung. 2018. GaAN: Gated Attention Networks for Learning on Large and Spatiotemporal Graphs. In *34th Conference on Uncertainty in Artificial Intelligence 2018, UAI 2018*.
- [24] Lingxiao Zhao and Leman Akoglu. 2019. PairNorm: Tackling Oversmoothing in GNNs. In *International Conference on Learning Representations*.