

Graph Convolutions Enrich the Self-Attention in Transformers!

Jeongwhan Choi*

Yonsei University

jeongwhan.choi@yonsei.ac.kr

Hyowon Wi*

KAIST

hyowon.wi@kaist.ac.kr

Jayoung Kim

KAIST

jayoung.kim@kaist.ac.kr

Yehjin Shin

KAIST

yehjin.shin@kaist.ac.kr

Kookjin Lee

Arizona State University

kookjin.lee@asu.edu

Nathaniel Trask

University of Pennsylvania

ntrask@seas.upenn.edu

Noseong Park[†]

KAIST

noseong@kaist.ac.kr

Abstract

Transformers, renowned for their self-attention mechanism, have achieved state-of-the-art performance across various tasks in natural language processing, computer vision, time-series modeling, etc. However, one of the challenges with deep Transformer models is the oversmoothing problem, where representations across layers converge to indistinguishable values, leading to significant performance degradation. We interpret the original self-attention as a simple graph filter and redesign it from a graph signal processing (GSP) perspective. We propose a graph-filter-based self-attention (GFSA)¹ to learn a general yet effective one, whose complexity, however, is slightly larger than that of the original self-attention mechanism. We demonstrate that GFSA improves the performance of Transformers in various fields, including computer vision, natural language processing, graph-level tasks, speech recognition, and code classification.

1 Introduction

Transformers are arguably one of the best feats in the field of deep learning. They are now showing state-of-the-art performance in various fields, ranging from computer vision to natural language processing, prediction tasks on graphs, speech recognition, and so forth [77, 16, 60, 61, 19, 74, 101, 45, 27, 41, 63, 51, 59, 38, 95, 72]. Recently, there have been several studies conducted on better understanding them [25, 3, 80]; there exists a common agreement among researchers that the self-attention is one of the keys leading to the success.

*Equal contribution.

[†]Corresponding author.

¹The source code of GFSA is available at: <https://github.com/jeongwhanchoi/GFSA>.

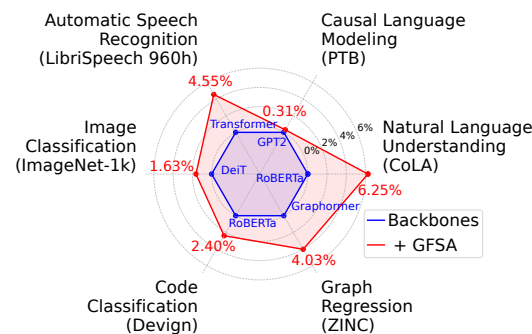


Figure 1: Performance improvements (%) of our GFSA when integrated with different Transformer backbones in various domains. We achieve these results with only tens to hundreds of additional parameters to Transformers.

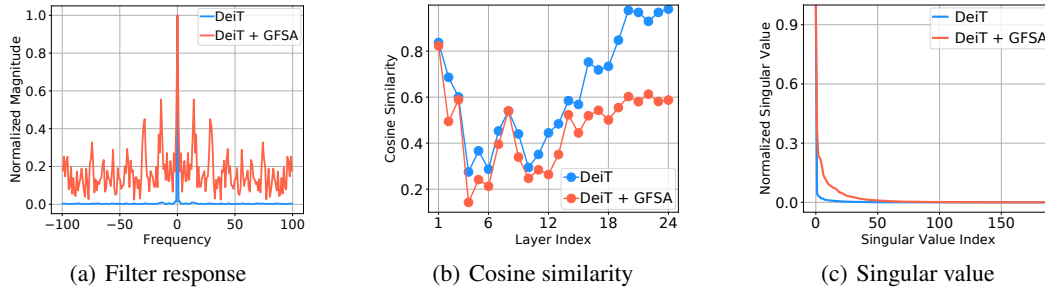


Figure 2: Filter frequency response, cosine similarity, and singular values on ImageNet-1k for DeiT-S and DeiT-S + GFSA. Details and more visualizations are in Appendices C and D.

However, there also exist several studies pointing out potential limitations of the self-attention [101, 18, 25, 28]. For instance, Shi et al. [71] revealed an analogy between the self-attention and the residual graph convolutional network (GCN), showing that BERT also suffers from a notorious problem of GCNs, called *oversmoothing*, i.e., tokens’ latent representations become similar to each other at the deeper layers. In every self-attention layer, value vectors are aggregated in a weighted average manner since each row-wise sum of the attention matrix is always 1. Although each self-attention layer has its own attention matrix, this aggregation method causes the oversmoothing problem, not only in Transformers but also in graph neural networks [54, 7, 79, 66, 37, 101, 25, 92, 53, 71, 80, 3, 89, 90]. However, we confine our discussion to the oversmoothing of Transformers (see Section 2).

Being inspired by them, we redesign the self-attention from the perspective of graph signal processing (GSP) — in particular, we resort to GSP on directed graphs since the attention matrix is asymmetric. However, performing graph convolutions in the self-attention layer may incur non-trivial computational overheads. Therefore, our key design point is to learn a general but effective graph filter with minimal overhead. In general, a graph filter on a graph $\mathcal{G}$ is represented by a polynomial expression based on its adjacency or Laplacian matrix — in this regard, the existing self-attention mechanism can be understood as the simplest graph filter with $\bar{\mathbf{A}}$ only, where $\bar{\mathbf{A}} \in [0, 1]^{n \times n}$ means a learned attention matrix and n is the number of input tokens.

Our proposed graph filter consists of an identity term and two matrix polynomial terms, $\bar{\mathbf{A}}$ and $\bar{\mathbf{A}}^K$. One can design better filters with more polynomial terms, but we avoid it since Transformers already require very large computation. The K -th power, $\bar{\mathbf{A}}^K$, may also require a high computation when the number of tokens is large. To avoid this, we further approximate $\bar{\mathbf{A}}^K$ using the element-wise first-order Taylor approximation. Therefore, one can consider that our proposed graph filter is the very next complicated filter after the one used by the original self-attention mechanism. However, its efficacy is tremendous in various fields (cf. Fig. 1).

Our proposed filter enriches the self-attention with more diverse frequency information (see Fig. 2(a)) — low (resp. high) frequency signals on $\mathcal{G}$ mean neighboring nodes have similar (resp. different) values. Therefore, our method is able to not only effectively address the oversmoothing problem but also learn better latent representations for downstream tasks.

There exist a couple of prior works to enrich the self-attention mechanism with high frequency information [80, 4]. In comparison with them, our proposed graph filter is distinctive in the following aspects: i) our proposed filter is more effective and shows better performance with comparable computational overheads, ii) our proposed filter is well-aligned with recent advancements in the GCN community — in other words, some graph filters used by recent advanced GCN methods are special cases of our proposed graph filter, which is not the case for prior works, and iii) other methods were typically studied for certain domains only whereas we test our method in 6 domains — for instance, DiversePatch [25] works only for Vision Transformers (ViTs).

We replace the self-attention layer of selected Transformers in various fields with our proposed graph filter-based layer without changing other parts. Therefore, the accuracy increases in them are solely by our proposed graph filter-based self-attention. These enriched Transformers increase the model performance by 1.63% for image classification, 6.25% for natural language understanding, 0.31% for

causal language modeling, 4.03% for graph regression, 4.76% for speech recognition, and 2.40% for code classification (see Fig. 1). Our core contributions are as follows:

- We provide a novel perspective on self-attention as a graph filter. This perspective allows us to design more effective self-attention that can address the oversmoothing problem.
- We propose a graph filter-based self-attention (GFSA) mechanism, integrating an identity term and two polynomial terms for general yet effective than the simple self-attention mechanism (Section 3).
- We demonstrate that GFSA improves the performance of Transformers on a variety of tasks. GFSA achieves improved results on natural language processing, computer vision, speech recognition, graph-level tasks, and code classification (Sections 5.1 to 5.6).
- We devise a strategy to selectively apply GFSA to even-numbered layers, effectively mitigating the computational overhead while preserving GFSA's performance (Section 6).

2 Background & Related Work

2.1 Self-Attention in Transformers

The core building block of the Transformer architecture is the self-attention mechanism, which enables the model to learn attention patterns over its input tokens [77]. The self-attention mechanism, denoted as $\text{SA} : \mathbb{R}^{n \times d} \rightarrow \mathbb{R}^{n \times d}$, can be expressed as follows:

$$\text{SA}(\mathbf{X}) = \text{softmax}\left(\frac{\mathbf{X}\mathbf{W}_{\text{qry}}(\mathbf{X}\mathbf{W}_{\text{key}})^\top}{\sqrt{d}}\right)\mathbf{X}\mathbf{W}_{\text{val}} = \bar{\mathbf{A}}\mathbf{X}\mathbf{W}_{\text{val}}, \quad (1)$$

where $\mathbf{X} \in \mathbb{R}^{n \times d}$ is the input feature matrix, $\mathbf{W}_{\text{key}} \in \mathbb{R}^{d \times d}$, $\mathbf{W}_{\text{qry}} \in \mathbb{R}^{d \times d}$, and $\mathbf{W}_{\text{val}} \in \mathbb{R}^{d \times d}$ are the key, query, and value trainable parameters, respectively, and d is the dimension of each token. The self-attention mechanism allows the model to weigh the importance of each token in the input sequence relative to the others, enabling the model to capture long-range contextual information better. The Transformer architecture includes multiple layers, each with a multi-head self-attention layer followed by a position-wise feed-forward layer.

2.2 Self-Attention and Graph Convolutional Filter

The self-attention matrix used in Transformers has the form of symmetrically normalized adjacency matrix where each token become a node [71, 28] — the symmetrically normalized adjacency matrix is a special case of asymmetric (or directed) adjacency matrix where each row is normalized and is frequently used for the graph signal processing (GSP) on directed graphs [49]. A weighted graph $\mathcal{G}$ with adjacency matrix $\mathbf{A}$ can be constructed by using the input tokens as n nodes and the edge weights between node i and node j as $\exp((\mathbf{X}\mathbf{W}_{\text{qry}})_i^\top (\mathbf{X}\mathbf{W}_{\text{key}})_j)$. We can rewrite the self-attention matrix $\bar{\mathbf{A}}_{ij}$ as $\frac{\exp((\mathbf{X}\mathbf{W}_{\text{qry}})_i^\top (\mathbf{X}\mathbf{W}_{\text{key}})_j)}{\sum_{k=1}^d \exp((\mathbf{X}\mathbf{W}_{\text{qry}})_i^\top (\mathbf{X}\mathbf{W}_{\text{key}})_k)}$. This allows $\bar{\mathbf{A}}$ to be interpreted as the symmetrically normalized adjacency matrix. In other words, $\bar{\mathbf{A}} = \mathbf{D}^{-1}\mathbf{A}$, where $\mathbf{D} = \text{diag}(d_1, d_2, \dots, d_n)$ and $d_i = \sum_j \mathbf{A}_{i,j}$.

Our new attention method is designed on top of GSP which has a close connection to discrete signal processing (DSP) [67, 68]. In DSP, a discrete signal with a length of n can be represented by a vector $\mathbf{x} \in \mathbb{R}^n$. Let $\mathbf{g} \in \mathbb{R}^n$ be a filter that we want to apply to $\mathbf{x}$. The convolution $\mathbf{x} * \mathbf{g}$ can be written as follows:

$$\mathbf{y}_i = \sum_{j=1}^n \mathbf{x}_j \mathbf{g}_{i-j}, \quad (2)$$

where the index, denoted as i , refers to the i -th element in each vector.

GSP can be understood as a generalized concept of DSP. Signals are defined on the nodes of a graph, and the graph's structure influences signal processing operations. In addition, the linear and shift-invariant graph convolution filter $\mathbf{H}$ with n nodes can be written with a shift operator $\mathbf{S}$ as

follows — $\mathbf{S}$ can be from a directed graph [48]:

$$\mathbf{y} = \mathbf{H}\mathbf{x} = \sum_{k=0}^K w_k \mathbf{S}^k \mathbf{x}, \quad (3)$$

where $\mathbf{x} \in \mathbb{R}^n$ is a 1-dimensional graph signal, K is the maximum order of polynomial, and $w_k \in [-\infty, \infty]$ is a coefficient. $\mathbf{S}$ is an $n \times n$ matrix where (i, j) -th element is non-zero if and only if there is an edge from node i to j . Two representative samples of $\mathbf{S}$ are adjacency and Laplacian matrices. The graph filter $\mathbf{H}$ is the same as $\sum_{k=0}^K w_k \mathbf{S}^k$ with a large enough value of K , which is called *matrix polynomial* [48]. We note that this graph filtering operation can be extended to d -dimensional cases as in Eq. (1). Being inspired by Zou et al. [106] and Maskey et al. [49], we rely on the singular value domain analysis to understand the low/high-pass characteristics of filters on directed graphs (cf. Fig. 2). See more discussion in Appendices E and F.

In the context of the self-attention within Transformers, the core part of the self-attention in Eq. (1), i.e., $\bar{\mathbf{A}}\mathbf{X}$, can be considered as a d -dimensional graph filter with $\bar{\mathbf{A}}$ only, where $\mathbf{H} = \bar{\mathbf{A}}$. Our goal in this paper is to design a simple (for computational purposes) yet effective form of $\mathbf{H}$.

2.3 Oversmoothing in GCNs and Transformers

Oversmoothing is a phenomenon observed in deep learning models, especially in GCNs [39, 78]. As information is aggregated over multiple layers for multiple nodes (tokens), latent representations tend to become similar to each other, leading to a loss of distinctiveness in the representations [54, 104, 66].

Surprisingly, an oversmoothing-like phenomenon is also observed in Transformers [80, 71]. Unlike CNNs, Transformers can not benefit from simply deepening layers after a certain depth. Earlier studies hypothesize that this may be due to issues such as the attention or feature collapse or due to uniformity among patches or tokens [101, 25, 92]. Dong et al. [18] also point out that the output of a pure Transformer, i.e., an attention mechanism without skip connections or MLPs, tends to converge to a rank-1 matrix [18]. This analysis is followed by [53], which suggests that rank collapses incur vanishing gradients of attention queries and keys.

In this context, the self-attention acts as a low-pass filter, since the self-attention calculates the weighted average of the value vectors of tokens. Wang et al. [80, Theorem 1] also reveal that the self-attention is a low-pass filter, continuously reducing high-frequency information. This nature contributes to the oversmoothing phenomenon as unique high-frequency features are lost in deeper layers of the network, further worsening the uniformity of token representations. Therefore, we extend the term “oversmoothing” to describe the degeneration challenge observed in Transformers.

There have been proposed many empirical countermeasures for ViT, such as patch diversification [102, 25], rank collapse alleviation [101, 99], and training stabilization [74, 98]. Similar alleviating methods have been also proposed in the field of NLP, such as unsupervised learning [9], and resolve the oversmoothing and the token uniformity (or information diffusion) problems [18, 92]. There are studies on utilizing high frequency information via frequency domain analyses [80, 4], but they are not designed on top of graph filtering perspectives. Dovonon et al. [20] find that Transformers are not inherently low-pass filters, but oversmoothing depends on the eigenspectrum of the self-attention layers. They propose a reparametrization of the Transformer weights, ensuring that oversmoothing does not occur.

Our paper addresses the oversmoothing problem with graph filters since the self-attention mechanism is a basic graph filtering operation as seen in the previous subsection.

3 Graph Filter-based Self-Attention Layers

Let $\bar{\mathbf{A}} \in [0, 1]^{n \times n}$, where n is the number of tokens in the input to the self-attention layer, be a self-attention matrix. Since Transformers use multi-head self-attentions, there are multiple such matrices. For simplicity, but without loss of generality, we discuss only one head in one layer.

From the GSP perspective, using $\bar{\mathbf{A}}$ as the shift operator, a graph filter can be represented as a matrix polynomial filter, as mentioned in Section 2.2. We aim to design this matrix polynomial filter using the two lowest-order terms and one high-order term in Eq. (3). The following theorem shows that,

despite using the three terms, the filter can be either a low-pass filter or a high-pass filter, depending on the coefficient values.

Theorem 3.1 (Filter characteristics based on coefficient values). *Let $\bar{\mathbf{A}}$ be a self-attention matrix interpreted as a graph with connected components. Consider the polynomial graph filter defined by $\sum_{k=0}^K w_k \bar{\mathbf{A}}^k$, where $w_2, w_3, \dots, w_{K-1} = 0$ and only w_0, w_1 , and w_K are non-zero. If the coefficients w_k for $k = 0, 1, K$ are positive and their sum is 1, then the polynomial filter acts as a low-pass filter, attenuating high-frequency components and promoting smoothness across the graph. Conversely, if $w_k = (-\alpha)^k$ for $k = 0, 1, K$ and $\alpha \in (0, 1)$ with sufficient large K , the polynomial filter exhibits high-pass filter behavior.*

The proof of Theorem 3.1 is in Appendix G. Based on Theorem 3.1, we propose to use the following graph filter, $\mathbf{H}_{\text{GFSA}}$, where the two lowest-order terms and one high-order term of the matrix polynomial are used:

$$\mathbf{H}_{\text{GFSA}} = w_0 \mathbf{I} + w_1 \bar{\mathbf{A}} + w_K \bar{\mathbf{A}}^K, \quad (4)$$

where w_0, w_1, w_K are coefficients and K is a hyper-parameter where $K \geq 2$. The coefficients can be learnable weights and we learn them with gradient descent algorithms.

Approximation of the high-order term. In Eq. (4), it is costly to calculate $\bar{\mathbf{A}}^K$ when K is large, so we need a way to approximate the high-order term $\bar{\mathbf{A}}^K$ in GFSA. We use the first-order Taylor approximation at point $a = 1$ for this purpose:

$$f(x) \simeq f(a) + f'(a)(x - a), \quad (5)$$

thus, we approximate $f(K) = \bar{\mathbf{A}}^K$ as follows:

$$f(K) = \bar{\mathbf{A}}^K \simeq f(1) + f'(1)(K - 1). \quad (6)$$

Computing the derivative of $\bar{\mathbf{A}}^K$ directly at the evaluation point requires high computational costs. To overcome this problem, we adopt the forward finite difference method, which approximates derivatives with the difference term:

$$f'(K) = \frac{f(K + h) - f(K)}{h} = \frac{\bar{\mathbf{A}}^{K+h} - \bar{\mathbf{A}}^K}{h}, \quad (7)$$

where the approximation error is $\mathcal{O}(h^2)$. To balance the trade-off between computational efficiency² and accuracy, we set $h = 1$. This method is inspired by the approach in Brouwer et al. [6], which uses the difference term between two consecutive hidden states in discrete time to approximate the derivatives. Therefore, we approximate $\bar{\mathbf{A}}^K$ as:

$$\begin{aligned} f(K) = \bar{\mathbf{A}}^K &\simeq f(1) + (\bar{\mathbf{A}}^2 - \bar{\mathbf{A}})(K - 1) \\ &= \bar{\mathbf{A}} + (K - 1)(\bar{\mathbf{A}}^2 - \bar{\mathbf{A}}). \end{aligned} \quad (8)$$

The approximation for $\bar{\mathbf{A}}^K$ with $\bar{\mathbf{A}}$ and $\bar{\mathbf{A}}^2$ provides a simpler computation that can significantly reduce the required computational resources and time.

GFSA: our graph filter-based self-attention. Our proposed graph filter-based self-attention (GFSA) is defined with the graph filter $\tilde{\mathbf{H}}_{\text{GFSA}}$ as follows:

$$\text{GFSA}(\mathbf{X}) := \tilde{\mathbf{H}}_{\text{GFSA}} \mathbf{X} \mathbf{W}_{\text{val}}, \quad (9)$$

$$\tilde{\mathbf{H}}_{\text{GFSA}} = w_0 \mathbf{I} + w_1 \bar{\mathbf{A}} + w_K (\bar{\mathbf{A}} + (K - 1)(\bar{\mathbf{A}}^2 - \bar{\mathbf{A}})), \quad (10)$$

where the last term is the approximated $\bar{\mathbf{A}}^K$ from Eq. (8). We replace the original self-attention layer in various Transformers with the proposed graph filter-based layer without changing other parts. Therefore, GFSA can be plugged into any Transformers that rely on the self-attention. For pseudocode, see Appendix I.

²Calculating the power of a matrix for small h requires a high computational cost since it is calculated by diagonalizing the matrix or using Schur normal form [31].

4 Properties of GFSA

This section analyzes the theoretical error of the $\bar{\mathbf{A}}^K$ approximation used by GFSA and how GFSA can mitigate oversmoothing. We also explain the meaning of GFSA's high-order term in the context of Transformers and provide comparisons of GFSA in other models.

Theoretical characteristics of approximation error in GFSA. We provide a theorem that provides an upper bound on the error introduced by approximating the power of a matrix, specifically using the first-order Taylor expansion. The following theorem specifically analyzes the error of matrix $\bar{\mathbf{A}}^K$ when approximated using a first-order Taylor expansion.

Theorem 4.1 (Error bound for approximated high-order term in GFSA). *Define the error term, E_K , as the difference between the exact value and approximated value of $\bar{\mathbf{A}}^K$, which is given by $E_K = \|\bar{\mathbf{A}}^K - (\bar{\mathbf{A}} + (K-1)(\bar{\mathbf{A}}^2 - \bar{\mathbf{A}}))\|_F$, where $\|\cdot\|_F$ denotes the Frobenius norm. Then, the error bound can be shown that $E_K \leq 2\sqrt{n}K$.*

The error bound provides an upper limit for the difference between the actual value of $\bar{\mathbf{A}}^K$ and its approximation. The proof of Theorem 4.1 is in Appendix H. It shows theoretical validity for using approximations to the high-order term in the filters of our GFSA. In terms of performance, we report a difference of approximately $\bar{\mathbf{A}}^K$ between the actual calculated values in Appendix J.

How to alleviate the oversmoothing problem? The key leading to the low/high pass filtering behavior of our proposed filter is the coefficients $\{w_0, w_1, w_K\}$ — note that in the self-attention of Transformers, $w_0 = w_K = 0$ and $w_1 = 1$. Since our method can learn any appropriate values for them for a downstream task, it can be reduced to low-pass-only, high-pass-only, or combined filters. According to Theorem 3.1, our graph polynomial filter can be said to be a low-pass filter when w_1, w_K are positive and a high-pass filter when they are negative. Therefore, our method can learn the appropriate coefficients $\{w_0, w_1, w_K\}$ for downstream tasks, so it can be reduced to a low-pass-only, high-pass-only, or combined filter, alleviating the oversmoothing problem of self-attention.

The meaning of the high-order term in GFSA in the context of Transformers. Existing self-attention only captures simple pairwise similarities between tokens and is limited in capturing high-order dependencies. For example, given the two sentences, “Books are more expensive than pencils” and “Books are cheaper than computers”, to understand the relationship between “computers” and “pencils”, we need to capture the high-order dependencies connected through the “Book” token. However, it is difficult to capture these high-order dependencies with traditional self-attention [103]. Therefore, from a Transformer perspective, the approximated $\bar{\mathbf{A}}^K$ in GFSA can be interpreted as being able to capture these high-order dependencies.

Comparison to Transformers. In the field of computer vision, there has been recent research on adjusting the frequency response of ViT. HAT [4] creates adversarial examples by altering clean images with high-frequency perturbations and jointly trains the ViT on clean images and adversarial examples. Through this, they aim to solve the problem of the ViT being unable to capture high-frequency by allowing us to capture the high-frequency components of the images. However, HAT has the disadvantage of requiring more epochs than the existing ViT, as it must perform adversarial training in some initial epochs and train normally in the remaining epochs. Wang et al. [80] use the concept of DSP, which is a special case of GSP, to isolate the lowest frequency component in the Fourier domain and use a filter learned by rescaling the low and high-frequency components. On the other hand, our GFSA extends the concept to graph signal processing and redesigns self-attention as a graph filter. While GFSA seeks to design a better graph filter by interpreting self-attention as a graph filter, Shi et al. [71] are inspired by JKNet [91], and they solve the oversmoothing problem by fusing the hidden vectors of each layer. However, their method has a limitation with memory increasing, and they only applied it to BERT.

Comparison to GCNs. Comparisons to GCNs that can be interpreted as graph filters [39, 15, 24] are inevitable. GFSA without a high-order term is analogous to ChebNet [15] with $K = 1$. In addition, GFSA reduces to the vanilla GCN [39] when $K = 1$, $w_0 = 0$, $w_1 = 1$. GPR-GNN [12], which approximates graph convolutions using the monomial basis, is identical to GFSA if it only considers up to first order and additionally uses a K -order term and learns the coefficients. When we

use only a high-order term and w_K is learned to a negative value, GFSA can become similar to the reaction-diffusion layer of GREAD [13], $\bar{A}X + \beta(\bar{A} - \bar{A}^2)$, depending on the higher order terms.

5 Experiments

In this section, we demonstrate the effectiveness of GFSA through a series of experiments. These experiments encompass various tasks: i) language understanding and causal language modeling, ii) image classification, iii) graph-level task, and iv) code classification. We replace the self-attention of base Transformers in those fields with our GFSA. Our modification adds only tens to hundreds of parameters, which are negligible in comparison with the original size of base models.

5.1 Experiments on Natural Language Understanding

Setting. We integrate GFSA into 3 pre-trained large language models: BERT, ALBERT, and RoBERTa. We evaluate them on the GLUE benchmark, which includes 3 categories of natural language understanding tasks: i) single-sentence, ii) similarity and paraphrasing, and iii) natural language inference tasks. For each task, we select the best hyperparameters for GFSA, and the other hyperparameters are fixed. The detailed experimental settings are in Appendix K.1.

Results. The results are shown in Table 1. When GFSA was plugged into backbones, average performance scores improved across all models over pure backbones. This indicates that GFSA is effective in both large models like BERT and RoBERTa, as well as relatively smaller models like ALBERT. It is worth mentioning that in the case of RoBERTa finetuned on the CoLA dataset; there is a significant margin increase from 60.34% to 64.11%, which is a 3.77% improvement with only 144 additional parameters. When compared to ContraNorm, GFSA shows a greater performance improvement on average. Fig. 5 in Appendix C shows that these performance enhancements can be attributed to addressing the oversmoothing issue through the designed graph filter.

Table 1: Results comparison on GLUE benchmark. **Avg** denotes the average performance.

Method	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	56.79	93.81	88.70	88.32	88.16	84.96/84.15	91.63	66.06	82.51
+ ContraNorm	110M	59.89	93.92	89.88	88.51	88.36	85.11/84.50	91.84	69.31	83.48
+ GFSA	110M	59.56	94.15	90.60	88.46	88.33	85.12/85.06	91.95	68.95	83.58
ALBERT _{BASE} [40]	11M	57.86	92.32	91.80	85.30	90.37	85.37/84.37	91.76	76.90	84.01
+ ContraNorm	11M	57.45	93.00	92.83	87.78	90.55	85.06/84.57	92.28	78.70	84.69
+ GFSA	11M	60.21	93.23	92.47	87.79	90.63	85.29/ 84.92	92.17	78.70	85.05
RoBERTa _{BASE} [44]	125M	60.34	94.84	92.28	88.86	89.99	87.94/87.30	92.57	78.70	85.87
+ ContraNorm	125M	63.06	95.41	93.17	88.91	90.34	87.88/87.40	92.82	80.51	86.61
+ GFSA	125M	64.11	95.41	93.52	89.09	90.35	87.99/87.54	92.97	80.14	86.79

5.2 Experiments on Causal Language Modeling

Setting. We also validate the effectiveness of GFSA on causal language modeling problems. We finetune GPT2 [61] on the following 3 datasets: Penn Treebank (PTB) [47], WikiText-2, and WikiText-103 [50]. Following the evaluation method in Yao et al. [93], we finetune models for 15 epochs with PTB, 4 epochs with WikiText-103, and 10 epochs with WikiText-2, and report the perplexity for sensitivity metric. The detailed experimental settings are in Appendix L.1.

Results. Table 2 shows the perplexity on PTB, WikiText-2, and WikiText-103. Across all datasets, GPT2 with GFSA consistently outperforms the vanilla GPT2. Our GFSA improves the average perplexity from 18.806 to 18.764. Note that performance improvements are made with only 144 additional learnable parameters for 12 layers with 12 heads.

Table 2: Results comparison on GPT-2 finetuned with GFSA. **Avg** denotes the average performance.

Method	#Params	PTB	WikiText-2	WikiText-103	Avg
GPT2	117M	19.513	20.966	15.939	18.806
+ GFSA	117M	19.450	20.923	15.919	18.764

5.3 Experiments on Vision Transformers

Setting. We aim to demonstrate the efficacy of our GFSA across a spectrum of ViT backbones. We choose DeiT [74], CaiT [75], and Swin [45] as the backbone, and the models are trained from scratch. When training the 12-layer DeiT, we follow the same training recipe, hyperparameters, and data augmentation from Touvron et al. [74]. For detailed experimental settings, see Appendix M.1.

Results. The experimental evaluations are summarized in Table 3. We compare various models on the ImageNet-1k benchmark. The results show that the proposed GFSA successfully enhances DeiT, CaiT, and Swin across all depth settings and training methods. GFSA provides additional parameters less than 72 for 12-layer DeiT while improving top-1 accuracy by 1.63%. To sum up, we observed that both shallow and deep ViTs can achieve the following benefits from GFSA: i) The filter response shows GFSA can preserve higher-frequency representation (cf. Fig. 2 (a)) and ii) Fig. 2 (b) shows that GFSA mitigates the increase in the cosine similarity of representation as the layer gets deeper. We further compare with state-of-the-art models that use Fourier transforms rather than graph filters in Appendix M.5. We also show results under the same settings as ContraNorm [28] in Appendix M.6.

Table 3: Results comparison on ImageNet-1k. Our full results with other models are in Appendix M.4.

Category	Method	Input Size	#Layers	#Params	Top-1 Acc
Transformer	DeiT-S [74]	224	12	22M	79.8
	DeiT-S + AttnScale [80]	224	12	22M	80.7
	DeiT-S + FeatScale [80]	224	12	22M	80.9
	DeiT-S + ContraNorm [28]	224	12	22M	80.4
	Swin-S [45]	224	12	50M	82.9
	DeiT-S [74]	224	24	43M	80.5
	CaiT-S [75]	224	24	47M	82.6
	DeiT-S + AttnScale [80]	224	24	44M	81.1
	DeiT-S + FeatScale [80]	224	24	44M	81.3
	DeiT-S + ContraNorm [28]	224	24	43M	80.7
GFSA	DeiT-S + GFSA	224	12	22M	81.1
	DeiT-S + GFSA	224	24	43M	81.5
	CaiT-S + GFSA	224	24	47M	82.8
	Swin-S + GFSA	224	12	50M	83.0

5.4 Experiments on Graph-level Tasks

Setting. To evaluate the efficacy of GFSA on graph-level tasks, we conduct experiments on a broader range of datasets. We use datasets from Long-Range Graph Benchmark (LRGB) [21] (e.g., Peptide-func and Peptide-struct), Benchmarking GNNs [22] (e.g., ZINC, MNIST, CIFAR10), Open Graph Benchmark (OGB) dataset [32] (e.g., Molhiv and MolTox21), and OGB-LSC dataset (i.e., PCQM4M-LSC) [33]. We choose Graphormer [94], Graph-ViT [30], and GPS [63] as our backbone architectures, following their original experimental protocols for fair comparison. For GPS, we replace its self-attention module with our GFSA while maintaining its best configuration and other hyperparameters. For Graph-ViT, we apply GFSA to the Hadamard self-attention method, which He et al. [30] propose as optimal. For a detailed experimental setting, see Appendix O.1.

Results. Tables 4, 5, and 6 show consistent performance improvements when GFSA is integrated with backbone architectures. Graph-ViT + GFSA shows improvements on all datasets. On Peptide-func, it achieves a 0.65% increase in AP. Notably, in PCQM4M, incorporating GFSA improves the validation MAE by 7.20%. Due to space constraints, the results with standard deviation are included in Appendix O.2.

5.5 Experiments on Automatic Speech Recognition

Setting. We conduct automatic speech recognition (ASR) experiments on the LibriSpeech ³ dataset [55], which consists of audio recordings paired with their transcriptions. We use Branch-

³<http://www.openslr.org/12>

Table 4: Results on ZINC			Table 5: Results on PCQM4M and PCQM4Mv2					
Method	#Params	MAE ($\downarrow$)	Method	#Params	PCQM4M		PCQM4Mv2	
					Train ($\downarrow$)	Validate ($\downarrow$)	Train ($\downarrow$)	Validate ($\downarrow$)
Graphormer	500K	0.1240	Graphormer	48.3M	0.0535	0.1286	0.0250	0.0862
+ GFSA	500K	0.1189	+ GFSA	48.3M	0.0312	0.1193	0.0249	0.0860

Table 6: Experimental evaluation of GFSA plugged into GPS and Graph-ViT. Results marked with $\dagger$ indicate settings where we conducted our own experiments due to unavailable Hadamard self-attention performance in He et al. [30]’s paper.

Method	Peptide-func	Peptide-struct	MNIST	CIFAR10	Molhiv	MolTOX21	ZINC
	AP ($\uparrow$)	MAE ($\downarrow$)	Accuracy ($\uparrow$)	Accuracy ($\uparrow$)	ROCAUC ($\uparrow$)	ROCAUC ($\uparrow$)	MAE ($\downarrow$)
GPS	0.6535 $\pm$ 0.0041	0.2500 $\pm$ 0.0005	0.9805 $\pm$ 0.0013	0.7230 $\pm$ 0.0036	–	–	0.070 $\pm$ 0.004
+ GFSA	0.6593 $\pm$ 0.0094	0.2496 $\pm$ 0.0013	0.9814 $\pm$ 0.0014	0.7244 $\pm$ 0.0048	–	–	0.069 $\pm$ 0.002
Graph-ViT	0.6919 $\pm$ 0.0085	$\dagger$ 0.2474 $\pm$ 0.0016	0.9820 $\pm$ 0.0005	0.6967 $\pm$ 0.0040	0.7792 $\pm$ 0.0149	0.7851 $\pm$ 0.0077	0.0849 $\pm$ 0.0047
+ GFSA	0.6964 $\pm$ 0.0025	0.2461 $\pm$ 0.0024	0.9826 $\pm$ 0.0004	0.6987 $\pm$ 0.0028	0.7830 $\pm$ 0.0109	0.7895 $\pm$ 0.0069	0.0845 $\pm$ 0.0032

former [59] and a pure Transformer. For implementation, we follow the recipes of SpeechBrain [65] and the detailed settings are in Appendix N.1.

Results. Table 7 compares word error rates (WERs) on LibriSpeech 100h and 960h. For 100h, Transformer+GFSA achieves 10.30/25.30 on the test clean/other set, which is a 6.53% improvement over the Transformer for the WER of the test clean. For 960h, Transformer+GFSA shows a WER result of 2.31 in test clean, a 4.55% improvement over Transformer and Branchformer+GFSA achieves 2.31/5.49 with an LM on the test clean/other sets. Fig. 8 in Appendix N.2 depicts the learning curves of train loss and valid loss when using GFSA, showing the effectiveness of our proposed filter.

Table 7: Results for ASR training on LibriSpeech 100h and 960h with GFSA

Method	#Params	LibriSpeech 100h		LibriSpeech 960h	
		test-clean WER	test-other WER	test-clean WER	test-other WER
Transformer	71.5M	11.02	25.42	2.42	5.50
+ GFSA	71.5M	10.30	24.30	2.31	5.49
Branchformer	109.8M	9.63	22.43	2.13	5.00
+ GFSA	109.8M	9.60	22.25	2.11	4.94

5.6 Experiments on Code Classification

Setting. We conduct a code defect detection task based on Devign dataset provided by Zhou et al. [105]. We use RoBERTa [44], CodeBERT [23], PLBART [2], and CodeT5 [84] as our backbone models. The detailed settings are in Appendix P.1.

Results. Table 8 shows the accuracy of all models; GFSA results better than the base models. The biggest improvement is 2.40% for RoBERTa. In the case of CodeT5-base, using GFSA shows an accuracy of 64.75, an improvement of 1.95% from 63.51. CodeT5-small+GFSA has only about 100 additional parameters compared to CodeT5-small with 60M parameters, and even more impressively, it surpasses the accuracy of CodeT5-base. The biggest improvement is 2.40% for RoBERTa. In Appendix P.2, we include case studies for this task. We also report the results of the code clone detection task in Appendix Q.

Table 8: Results on code classification. The number in ($\uparrow$) indicates the improvement rate.

Method	Accuracy
RoBERTa	62.88
+ GFSA	64.39 ($\uparrow$ 2.40%)
CodeBERT	64.31
+ GFSA	64.49 ($\uparrow$ 0.12%)
PLBART	62.63
+ GFSA	62.96 ($\uparrow$ 0.52%)
CodeT5-small	63.25
+ GFSA	63.69 ($\uparrow$ 0.70%)
CodeT5-base	63.51
+ GFSA	64.75 ($\uparrow$ 1.95%)

6 Discussion on Runtime Overheads

Limitation. The introduction of our GFSA layer results in a slight increase in training and inference time. We report the runtimes when plugging GFSA in Appendices R and S. For GLUE benchmark, integrating GFSA into BERT enhances the average performance from 82.51% to 83.58% (see Table 1) with more overhead of less than 36 seconds per epoch based on average training time (see Table 23). Considering the improvements, the increases in training time are negligible.

GFSA in selected layers: a strategy to mitigate the limitation. As GFSA requires more calculation than the original self-attention, the runtime after using GFSA slightly increases. Our experiments initially applied GFSA across all Transformer layers (as discussed in Section 5); however, to reduce computational load, we propose a selective application strategy. For this purpose, GFSA is used only on even-numbered layers. In Tables 35 to 39 of Appendix T, the results show that this strategy effectively reduces runtime increases while preserving comparable performance to the full-layer GFSA integration. Notably, the selective application of GFSA cuts the per-epoch runtime increase by 26.90% relative to its full-layer application, with only a 7.39% increase in runtime per epoch compared to the backbone model in Table 39.

GFSA in linear Transformers. Although GFSA requires additional computation for calculating $\tilde{\mathbf{A}}^2$, we explore integrating GFSA with linear attention variants to maintain efficiency and scalability. Recent approaches [36, 70] achieve linear complexity by reformulating softmax operations and reordering matrix multiplication in self-attention. We apply similar principles to compute second-order self-attention efficiently, enabling $\tilde{\mathbf{H}}_{\text{GFSA}}$ calculation with linear complexity with respect to sequence length. Fig. 4 shows the performance, runtime and GPU usage changes when applying our GFSA to Transformers with linear complexity. GFSA still improves performance compared to the backbone model, while the increase in time and GPU usage is minimal. Notably, when GFSA is applied to Efficient Attention [70], the performance is improved and the runtime is 11.82 times faster than when GFSA is applied to the vanilla self-attention. This shows that GFSA can be effectively implemented with linear complexity architectures while preserving its benefits and providing a solution for addressing computational concerns.

7 Conclusion

Our proposed GFSA achieves high performance with improvements on a variety of tasks. GFSA is a simple yet effective method that enriches self-attention in Transformers with more diverse frequency information. This enables GFSA to address the oversmoothing problem and learn better latent representations for downstream tasks. However, our GFSA does not bring significant overheads in those Transformers' empirical runtime complexity. One can use more complicated graph filters to enhance accuracy more, but our goal is to find a balance between accuracy enhancements and overheads in runtime complexity.

We believe that GFSA suggests a promising new direction for improving Transformers. GFSA can be implemented with a simple way and used in conjunction with other techniques to further improve the performance of Transformers. Considering the ongoing advancements in large language models, such as GPT-4 [1] and LLaMA [76], we hope that our approach may offer new insights for enhancing their performance and efficiency.

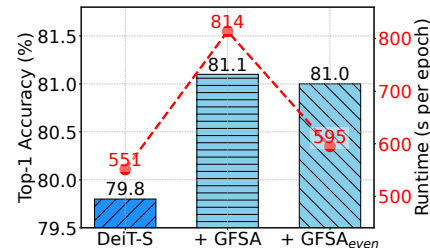


Figure 3: Effectiveness of our selective layer strategy on ImageNet-1k. This shows our strategy's ability to maintain accuracy benefits while mitigating runtime increases.

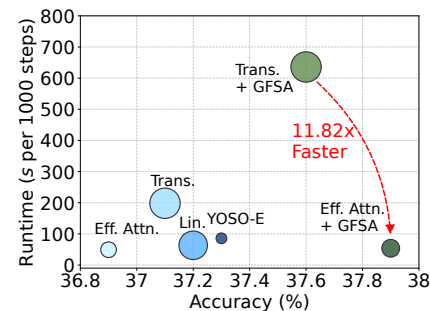


Figure 4: Performance (x -axis), runtime (y -axis), and GPU usage (circle sizes) of various Transformers and integrated GFSA on Long-Range benchmark

Acknowledgements

N. Park was partly supported by the Korea Advanced Institute of Science and Technology (KAIST) grant funded by the Korea government (MSIT) (No. G04240001, Physics-inspired Deep Learning, 10%), Institute for Information & Communications Technology Planning & Evaluation (IITP) grants funded by the Korea government (MSIT) (No. RS-2020-II201361, Artificial Intelligence Graduate School Program (Yonsei University), 20%; No. RS-2024-00457882, AI Research Hub Project, 50%), and Samsung Electronics Co., Ltd. (No. G01240136, KAIST Semiconductor Research Fund (2nd), 10%). K. Lee acknowledges support from the U.S. National Science Foundation under grant IIS 2338909. Dr. Trask acknowledges funding under the Department of Energy under the Mathematical Multifaceted Integrated Capability Centers program.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altmenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [2] Wasi Ahmad, Saikat Chakraborty, Baishakhi Ray, and Kai-Wei Chang. Unified pre-training for program understanding and generation. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 2655–2668, 2021.
- [3] Ameen Ali, Tomer Galanti, and Lior Wolf. Centered self-attention layers. *arXiv preprint arXiv: 2306.01610*, 2023.
- [4] Jiawang Bai, Li Yuan, Shu-Tao Xia, Shuicheng Yan, Zhifeng Li, and Wei Liu. Improving vision transformers by revisiting high-frequency components. In *European Conference on Computer Vision*, pages 1–18. Springer, 2022.
- [5] Luisa Bentivogli, Peter Clark, Ido Dagan, and Danilo Giampiccolo. The fifth pascal recognizing textual entailment challenge. *TAC*, 7:8, 2009.
- [6] Edward De Brouwer, Jaak Simm, Adam Arany, and Yves Moreau. GRU-ODE-Bayes: Continuous modeling of sporadically-observed time series. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019.
- [7] Chen Cai and Yusu Wang. A note on over-smoothing for graph neural networks. *arXiv preprint arXiv:2006.13318*, 2020.
- [8] Daniel Cer, Mona Diab, Eneko Agirre, Inigo Lopez-Gazpio, and Lucia Specia. Semeval-2017 task 1: Semantic textual similarity-multilingual and cross-lingual focused evaluation. *arXiv preprint arXiv:1708.00055*, 2017.
- [9] Nuo Chen, Linjun Shou, Ming Gong, Jian Pei, Bowen Cao, Jianhui Chang, Daxin Jiang, and Jia Li. Alleviating over-smoothing for unsupervised sentence representation. *arXiv preprint arXiv:2305.06154*, 2023.
- [10] Tianlong Chen, Zhenyu Zhang, Yu Cheng, Ahmed Awadallah, and Zhangyang Wang. The principle of diversity: Training stronger vision transformers calls for reducing all levels of redundancy. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12020–12030, 2022.
- [11] Zihan Chen, Hongbo Zhang, Xiaoji Zhang, and Leqi Zhao. Quora question pairs, 2018.
- [12] Eli Chien, Jianhao Peng, Pan Li, and Olga Milenkovic. Adaptive universal generalized PageRank graph neural network. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2021.
- [13] Jeongwhan Choi, Seoyoung Hong, Noseong Park, and Sung-Bae Cho. Gread: Graph neural reaction-diffusion networks. In *International Conference on Machine Learning (ICML)*, pages 5722–5747. PMLR, 2023.

- [14] George Dasoulas, Kevin Scaman, and Aladin Virmaux. Lipschitz normalization for self-attention layers with application to graph neural networks. In *International Conference on Machine Learning (ICML)*, pages 2456–2466. PMLR, 2021.
- [15] Michaël Defferrard, Xavier Bresson, and Pierre Vandergheynst. Convolutional neural networks on graphs with fast localized spectral filtering. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2016.
- [16] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1423. URL <https://aclanthology.org/N19-1423>.
- [17] Bill Dolan and Chris Brockett. Automatically constructing a corpus of sentential paraphrases. In *Third International Workshop on Paraphrasing (IWP2005)*, 2005.
- [18] Yihe Dong, Jean-Baptiste Cordonnier, and Andreas Loukas. Attention is not all you need: Pure attention loses rank doubly exponentially with depth. In *International Conference on Machine Learning (ICML)*, pages 2793–2803. PMLR, 2021.
- [19] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2021.
- [20] Gbètondji JS Dovonon, Michael M Bronstein, and Matt J Kusner. Setting the record straight on transformer oversmoothing. *arXiv preprint arXiv:2401.04301*, 2024.
- [21] Vijay Prakash Dwivedi, Ladislav Rampášek, Mikhail Galkin, Ali Parviz, Guy Wolf, Anh Tuan Luu, and Dominique Beaini. Long range graph benchmark. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022. URL <https://openreview.net/forum?id=in7XC5RcjEn>.
- [22] Vijay Prakash Dwivedi, Chaitanya K Joshi, Anh Tuan Luu, Thomas Laurent, Yoshua Bengio, and Xavier Bresson. Benchmarking graph neural networks. *Journal of Machine Learning Research*, 24(43):1–48, 2023.
- [23] Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xiaocheng Feng, Ming Gong, Linjun Shou, Bing Qin, Ting Liu, Daxin Jiang, et al. CodeBERT: A pre-trained model for programming and natural languages. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 1536–1547, 2020.
- [24] Johannes Gasteiger, Stefan Weißenberger, and Stephan Günnemann. Diffusion improves graph learning. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 32, 2019.
- [25] Chengyue Gong, Dilin Wang, Meng Li, Vikas Chandra, and Qiang Liu. Vision transformers with patch diversification. *arXiv preprint arXiv:2104.12753*, 2021.
- [26] Alex Graves and Navdeep Jaitly. Towards end-to-end speech recognition with recurrent neural networks. In *International Conference on Machine Learning (ICML)*, pages 1764–1772. PMLR, 2014.
- [27] Anmol Gulati, James Qin, Chung-Cheng Chiu, Niki Parmar, Yu Zhang, Jiahui Yu, Wei Han, Shibo Wang, Zhengdong Zhang, Yonghui Wu, et al. Conformer: Convolution-augmented transformer for speech recognition. *arXiv preprint arXiv:2005.08100*, 2020.
- [28] Xiaojun Guo, Yifei Wang, Tianqi Du, and Yisen Wang. ContraNorm: A contrastive learning perspective on oversmoothing and beyond. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2023.

- [29] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [30] Xiaoxin He, Bryan Hooi, Thomas Laurent, Adam Perold, Yann LeCun, and Xavier Bresson. A generalization of vit/mlp-mixer to graphs. In *International conference on machine learning (ICML)*, pages 12724–12745. PMLR, 2023.
- [31] Nicholas J Higham. *Functions of matrices: theory and computation*. SIAM, 2008.
- [32] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. Open graph benchmark: Datasets for machine learning on graphs. *Advances in neural information processing systems (NeurIPS)*, 33:22118–22133, 2020.
- [33] Weihua Hu, Matthias Fey, Hongyu Ren, Maho Nakata, Yuxiao Dong, and Jure Leskovec. OGB-LSC: A large-scale challenge for machine learning on graphs. *arXiv preprint arXiv:2103.09430*, 2021.
- [34] Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely connected convolutional networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4700–4708, 2017.
- [35] John J Irwin, Teague Sterling, Michael M Mysinger, Erin S Bolstad, and Ryan G Coleman. Zinc: a free tool to discover chemistry for biology. *Journal of chemical information and modeling*, 52(7):1757–1768, 2012.
- [36] Angelos Katharopoulos, Apoorv Vyas, Nikolaos Pappas, and François Fleuret. Transformers are rnns: Fast autoregressive transformers with linear attention. In *International conference on machine learning*, pages 5156–5165. PMLR, 2020.
- [37] Nicolas Keriven. Not too little, not too much: a theoretical analysis of graph (over) smoothing. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 2268–2281, 2022.
- [38] Jayoung Kim, Yehjin Shin, Jeongwhan Choi, Hyowon Wi, and Noseong Park. Polynomial-based self-attention for table representation learning. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 24509–24526. PMLR, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/kim24ae.html>.
- [39] Thomas N. Kipf and Max Welling. Semi-supervised classification with graph convolutional networks. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2017.
- [40] Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. Albert: A lite bert for self-supervised learning of language representations. *arXiv preprint arXiv:1909.11942*, 2019.
- [41] Siddique Latif, Aun Zaidi, Heriberto Cuayahuitl, Fahad Shamshad, Moazzam Shoukat, and Junaid Qadir. Transformers in speech processing: A survey. *arXiv preprint arXiv:2303.11607*, 2023.
- [42] James Lee-Thorp, Joshua Ainslie, Ilya Eckstein, and Santiago Ontanon. Fnet: Mixing tokens with fourier transforms. *arXiv preprint arXiv:2105.03824*, 2021.
- [43] Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. BART: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7871–7880, 2020.

- [44] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. RoBERTa: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*, 2019.
- [45] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 10012–10022, 2021.
- [46] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*, 2017.
- [47] Mitchell Marcus, Beatrice Santorini, and Mary Ann Marcinkiewicz. Building a large annotated corpus of english: The penn treebank. 1993.
- [48] Antonio G. Marques, Santiago Segarra, and Gonzalo Mateos. Signal processing on directed graphs: The role of edge directionality when processing and learning from network data. *IEEE Signal Processing Magazine*, 37(6):99–116, 2020. doi: 10.1109/MSP.2020.3014597.
- [49] Sohir Maskey, Raffaele Paolino, Aras Bacho, and Gitta Kutyniok. A fractional graph laplacian approach to oversmoothing. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- [50] Stephen Merity, Caiming Xiong, James Bradbury, and Richard Socher. Pointer sentinel mixture models. *arXiv preprint arXiv:1609.07843*, 2016.
- [51] Luis Müller, Mikhail Galkin, Christopher Morris, and Ladislav Rampásek. Attending to graph transformers. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856. URL <https://openreview.net/forum?id=HhbqHBBrfZ>.
- [52] Tam Nguyen, Tan Nguyen, and Richard Baraniuk. Mitigating over-smoothing in transformers via regularized nonlocal functionals. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36, 2023.
- [53] Lorenzo Noci, Sotiris Anagnostidis, Luca Biggio, Antonio Orvieto, Sidak Pal Singh, and Aurelien Lucchi. Signal propagation in transformers: Theoretical perspectives and the role of rank collapse. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 27198–27211, 2022.
- [54] Kenta Oono and Taiji Suzuki. Graph neural networks exponentially lose expressive power for node classification. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2020.
- [55] Vassil Panayotov, Guoguo Chen, Daniel Povey, and Sanjeev Khudanpur. Librispeech: an asr corpus based on public domain audio books. In *2015 IEEE international conference on acoustics, speech and signal processing (ICASSP)*, pages 5206–5210. IEEE, 2015.
- [56] Daniel S Park, William Chan, Yu Zhang, Chung-Cheng Chiu, Barret Zoph, Ekin D Cubuk, and Quoc V Le. SpecAugment: A simple data augmentation method for automatic speech recognition. *Interspeech 2019*, 2019.
- [57] Badri Patro and Vijay Agneeswaran. Scattering vision transformer: Spectral mixing matters. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36, 2023.
- [58] Badri N Patro, Vinay P Namboodiri, and Vijay Srinivas Agneeswaran. Spectformer: Frequency and attention is what you need in a vision transformer. *arXiv preprint arXiv:2304.06446*, 2023.
- [59] Yifan Peng, Siddharth Dalmia, Ian Lane, and Shinji Watanabe. Branchformer: Parallel mlp-attention architectures to capture local and global context for speech recognition and understanding. In *International Conference on Machine Learning (ICML)*, pages 17627–17643. PMLR, 2022.
- [60] Alec Radford, Karthik Narasimhan, Tim Salimans, Ilya Sutskever, et al. Improving language understanding by generative pre-training. 2018.

- [61] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- [62] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *The Journal of Machine Learning Research*, 21(1):5485–5551, 2020.
- [63] Ladislav Rampášek, Michael Galkin, Vijay Prakash Dwivedi, Anh Tuan Luu, Guy Wolf, and Dominique Beaini. Recipe for a general, powerful, scalable graph transformer. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 14501–14515, 2022.
- [64] Yongming Rao, Wenliang Zhao, Zheng Zhu, Jiwen Lu, and Jie Zhou. Global filter networks for image classification. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 34, pages 980–993, 2021.
- [65] Mirco Ravanelli, Titouan Parcollet, Peter Plantinga, Aku Rouhe, Samuele Cornell, Loren Lugosch, Cem Subakan, Nauman Dawlatabad, Abdelwahab Heba, Jianyuan Zhong, Ju-Chieh Chou, Sung-Lin Yeh, Szu-Wei Fu, Chien-Feng Liao, Elena Rastorgueva, François Grondin, William Aris, Hwidong Na, Yan Gao, Renato De Mori, and Yoshua Bengio. SpeechBrain: A general-purpose speech toolkit, 2021. arXiv:2106.04624.
- [66] T. Konstantin Rusch, Michael M. Bronstein, and Siddhartha Mishra. A survey on oversmoothing in graph neural networks. *arXiv preprint arXiv: Arxiv-2303.10993*, 2023.
- [67] Aliaksei Sandryhaila and José MF Moura. Discrete signal processing on graphs. *IEEE transactions on signal processing*, 61(7):1644–1656, 2013.
- [68] Aliaksei Sandryhaila and Jose MF Moura. Discrete signal processing on graphs: Frequency analysis. *IEEE Transactions on Signal Processing*, 62(12):3042–3054, 2014.
- [69] Roy Schwartz, Jesse Dodge, Noah A Smith, and Oren Etzioni. Green AI. *Communications of the ACM*, 63(12):54–63, 2020.
- [70] Zhuoran Shen, Mingyuan Zhang, Haiyu Zhao, Shuai Yi, and Hongsheng Li. Efficient attention: Attention with linear complexities. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, pages 3531–3539, 2021.
- [71] Han Shi, Jiahui Gao, Hang Xu, Xiaodan Liang, Zhenguo Li, Lingpeng Kong, Stephen Lee, and James T Kwok. Revisiting over-smoothing in bert from the perspective of graph. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2022.
- [72] Yehjin Shin, Jeongwhan Choi, Hyowon Wi, and Noseong Park. An attentive inductive bias for sequential recommendation beyond the self-attention. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 8984–8992, 2024.
- [73] Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D Manning, Andrew Y Ng, and Christopher Potts. Recursive deep models for semantic compositionality over a sentiment treebank. In *Proceedings of the 2013 conference on empirical methods in natural language processing*, pages 1631–1642, 2013.
- [74] Hugo Touvron, Matthieu Cord, Matthijs Douze, Francisco Massa, Alexandre Sablayrolles, and Hervé Jégou. Training data-efficient image transformers & distillation through attention. In *International Conference on Machine Learning (ICML)*, pages 10347–10357. PMLR, 2021.
- [75] Hugo Touvron, Matthieu Cord, Alexandre Sablayrolles, Gabriel Synnaeve, and Hervé Jégou. Going deeper with image transformers. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 32–42, 2021.
- [76] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.
- [77] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in neural information processing systems (NeurIPS)*, volume 30, 2017.

- [78] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph Attention Networks. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2018.
- [79] Guangtao Wang, Rex Ying, Jing Huang, and Jure Leskovec. Multi-hop attention graph neural network. In *IJCAI*, 2021.
- [80] Peihao Wang, Wenqing Zheng, Tianlong Chen, and Zhangyang Wang. Anti-oversmoothing in deep vision transformers via the fourier domain analysis: From theory to practice. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2022.
- [81] Sinong Wang, Belinda Z Li, Madian Khabsa, Han Fang, and Hao Ma. Linformer: Self-attention with linear complexity. *arXiv preprint arXiv:2006.04768*, 2020.
- [82] Wei Wang, Ming Yan, and Chen Wu. Multi-granularity hierarchical attention fusion networks for reading comprehension and question answering. *arXiv preprint arXiv:1811.11934*, 2018.
- [83] Wenhan Wang, Ge Li, Bo Ma, Xin Xia, and Zhi Jin. Detecting code clones with graph neural network and flow-augmented abstract syntax tree. In *2020 IEEE 27th International Conference on Software Analysis, Evolution and Reengineering (SANER)*, pages 261–271, 2020. doi: 10.1109/SANER48275.2020.9054857.
- [84] Yue Wang, Weishi Wang, Shafiq Joty, and Steven CH Hoi. CodeT5: Identifier-aware unified pre-trained encoder-decoder models for code understanding and generation. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 8696–8708, 2021.
- [85] Alex Warstadt, Amanpreet Singh, and Samuel R Bowman. Neural network acceptability judgments. *Transactions of the Association for Computational Linguistics*, 7:625–641, 2019.
- [86] Ross Wightman. Pytorch image models. <https://github.com/rwightman/pytorch-image-models>, 2019.
- [87] Adina Williams, Nikita Nangia, and Samuel R Bowman. A broad-coverage challenge corpus for sentence understanding through inference. *arXiv preprint arXiv:1704.05426*, 2017.
- [88] Haiping Wu, Bin Xiao, Noel Codella, Mengchen Liu, Xiyang Dai, Lu Yuan, and Lei Zhang. CvT: Introducing convolutions to vision transformers. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 22–31, 2021.
- [89] Xinyi Wu, Amir Ajorlou, Zihui Wu, and Ali Jadbabaie. Demystifying oversmoothing in attention-based graph neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
- [90] Xinyi Wu, Zhengdao Chen, William Wei Wang, and Ali Jadbabaie. A non-asymptotic analysis of oversmoothing in graph neural networks. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2023.
- [91] Keyulu Xu, Chengtao Li, Yonglong Tian, Tomohiro Sonobe, Ken-ichi Kawarabayashi, and Stefanie Jegelka. Representation learning on graphs with jumping knowledge networks. In *International Conference on Machine Learning (ICML)*, pages 5453–5462, 2018.
- [92] Hanqi Yan, Lin Gui, Wenjie Li, and Yulan He. Addressing token uniformity in transformers via singular value transformation. In *Uncertainty in Artificial Intelligence*, pages 2181–2191. PMLR, 2022.
- [93] Zhewei Yao, Xiaoxia Wu, Conglong Li, Connor Holmes, Minjia Zhang, Cheng Li, and Yuxiong He. Random-ltd: Random and layerwise token dropping brings efficient training for large-scale transformers. *arXiv preprint arXiv:2211.11586*, 2022.
- [94] Chengxuan Ying, Tianle Cai, Shengjie Luo, Shuxin Zheng, Guolin Ke, Di He, Yanming Shen, and Tie-Yan Liu. Do transformers really perform badly for graph representation? In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 34, pages 28877–28888, 2021.

- [95] Youn-Yeol Yu, Jeongwhan Choi, Woojin Cho, Kookjin Lee, Nayong Kim, Kiseok Chang, ChangSeung Woo, ILHO KIM, SeokWoo Lee, Joon Young Yang, SOOYOUNG YOON, and Noseong Park. Learning flexible body collision dynamics with hierarchical contact mesh transformer. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=90yw2uM6J5>.
- [96] Li Yuan, Yunpeng Chen, Tao Wang, Weihao Yu, Yujun Shi, Zi-Hang Jiang, Francis EH Tay, Jiashi Feng, and Shuicheng Yan. Tokens-to-token vit: Training vision transformers from scratch on imagenet. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 558–567, 2021.
- [97] Zhanpeng Zeng, Yunyang Xiong, Sathya Ravi, Shailesh Acharya, Glenn M Fung, and Vikas Singh. You only sample (almost) once: Linear cost self-attention via bernoulli sampling. In *International conference on machine learning*, pages 12321–12332. PMLR, 2021.
- [98] Shuangfei Zhai, Tatiana Likhomanenko, Etai Littwin, Dan Busbridge, Jason Ramapuram, Yizhe Zhang, Jiatao Gu, and Joshua M. Susskind. Stabilizing transformer training by preventing attention entropy collapse. In *Proceedings of the 40th International Conference on Machine Learning (ICML)*, volume 202, pages 40770–40803. PMLR, 2023.
- [99] Aston Zhang, Alvin Chan, Yi Tay, Jie Fu, Shuohang Wang, Shuai Zhang, Huajie Shao, Shuochao Yao, and Roy Ka-Wei Lee. On orthogonality constraints for transformers. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)*, pages 375–382, 2021.
- [100] Ying Zhang, Mohammad Pezeshki, Philémon Brakel, Saizheng Zhang, César Laurent, Yoshua Bengio, and Aaron Courville. Towards end-to-end speech recognition with deep convolutional neural networks. *Interspeech 2016*, 2016.
- [101] Daquan Zhou, Bingyi Kang, Xiaojie Jin, Linjie Yang, Xiaochen Lian, Zihang Jiang, Qibin Hou, and Jiashi Feng. DeepViT: Towards deeper vision transformer. *arXiv preprint arXiv:2103.11886*, 2021.
- [102] Daquan Zhou, Yujun Shi, Bingyi Kang, Weihao Yu, Zihang Jiang, Yuan Li, Xiaojie Jin, Qibin Hou, and Jiashi Feng. Refiner: Refining self-attention for vision transformers. *arXiv preprint arXiv:2106.03714*, 2021.
- [103] Haoyi Zhou, Siyang Xiao, Shanghang Zhang, Jieqi Peng, Shuai Zhang, and Jianxin Li. Jump self-attention: Capturing high-order statistics in transformers. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 17899–17910, 2022.
- [104] Jie Zhou, Ganqu Cui, Shengding Hu, Zhengyan Zhang, Cheng Yang, Zhiyuan Liu, Lifeng Wang, Changcheng Li, and Maosong Sun. Graph neural networks: A review of methods and applications. *AI open*, 1:57–81, 2020.
- [105] Yaqin Zhou, Shangqing Liu, Jingkai Siow, Xiaoning Du, and Yang Liu. Devign: Effective vulnerability identification by learning comprehensive program semantics via graph neural networks. In *Advances in neural information processing systems (NeurIPS)*, volume 32, 2019.
- [106] Chunya Zou, Andi Han, Lequan Lin, and Junbin Gao. A simple yet effective svd-gcn for directed graphs. *arXiv preprint arXiv:2205.09335*, 2022.

Appendix

A Reproducibility Statement

To ensure the reproducibility and completeness of this paper, we include the Appendix with 12 sections. Appendix I provides our PyTorch-style pseudo code for our GFSA method. The pseudo code helps to implement our GFSA to any Transformers used a pure self-attention. All experiments in the paper are reproducible with additional implementation details provided in Appendices K to Q.

B Broader Impact

In terms of the broader impact of this research on society, we do not see the very negative impacts that might be expected. However, this paper may have implications for the carbon footprint and accessibility of learning algorithms. The computations required for machine learning research are rapidly growing, resulting in a larger carbon footprint [69]. Our study improves performance and increases runtime very slightly, but the runtime increase is not very significant. However, in future research, it will also be important to study and improve our GFSA by taking carbon footprints into account.

GFSA improves the performance of existing Transformer-based models, which can have many positive impacts on society through services that utilize natural language processing, computer vision, and speech recognition. However, it will also be important to improve GFSA by considering other dimensions of AI, such as robustness to adversarial examples, fairness, and explainability.

C Oversmoothing and Additional Visualizations

In Fig. 2, we show the visualizations of oversmoothing characteristics in DeiT. We also provide visualizations in other domains. We show the filter response, cosine similarity, and singular value of BERT finetuned on STS-B dataset of GLUE tasks in Fig. 5 and Graphormer finetuned on ZINC dataset in Fig. 6.

To characterize self-attention, we first analyze the filter response of self-attention in the frequency domain. We follow the method used by Wang et al. [80] for spectral visualization of the self-attention matrix. As shown in Fig. 2 (a), DeiT has a near-zero magnitude for the high frequencies, which is characteristic of a low-frequency filter and is likely to result in oversmoothing when applied multiple times.

We follow the calculation method of Guo et al. [28] for cosine similarity. As shown in Fig. 2 (b), the higher similarity as the layers of the model get deeper is related to the oversmoothing problem. To further analyze this issue, we also consider the dimensionality collapse in Transformer-based models. We plot the singular value distribution of the feature in the last block. As shown in Fig. 2 (c), insignificant, near-zero values dominate the feature distribution. As layers get deeper, the similarity of features increases and dimensional collapse occurs. The oversmoothing problem is the same in BERT and Graphormer, as shown in Fig. 5 and Fig. 6.

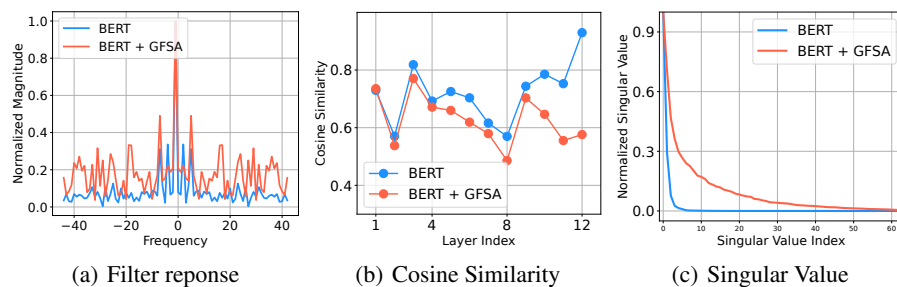


Figure 5: Filter frequency response, cosine similarity, and singular values on STS-B for BERT and BERT+GFSA

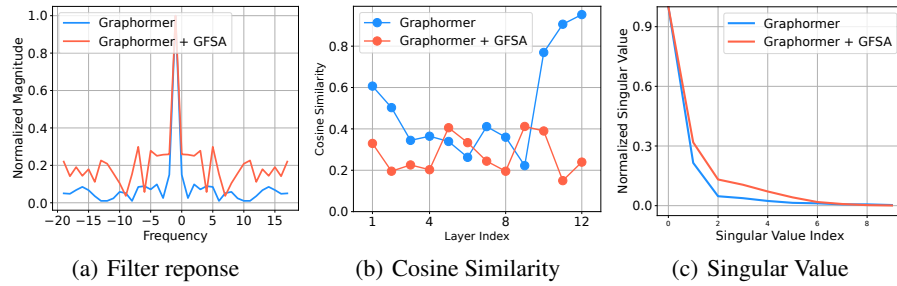


Figure 6: Filter frequency response, cosine similarity, and singular values on ZINC for Graphormer and Graphormer+GFSA

D Analysis of Frequency Responses with Visualizations

We analyze the frequency responses, which represent the impact of learned coefficients, for all 12 layers of BERT_{BASE} with and without GFSA. From Fig. 7, our analysis reveals that GFSA learns various filter types between layers. In early layers, we observe a tendency towards low-pass filtering, with prominent peaks at low frequencies. This aligns with the need for broader feature extraction in initial layers. The middle layers show a mix of low-pass and high-pass characteristics, with more complex frequency responses. This suggests GFSA is learning to balance between feature extraction and refinement. In deep layers, there is a noticeable shift towards higher frequency responses, indicating a move towards high-pass filtering. This shift supports our claim that GFSA can mitigate oversmoothing in deeper layers. BERT_{BASE}+GFSA shows a consistently higher magnitude response at higher frequencies, especially in deeper layers, compared to vanilla BERT. In other word, vanilla self-attention works primarily as a low-pass filter, while GFSA utilizes a wider range of frequencies.

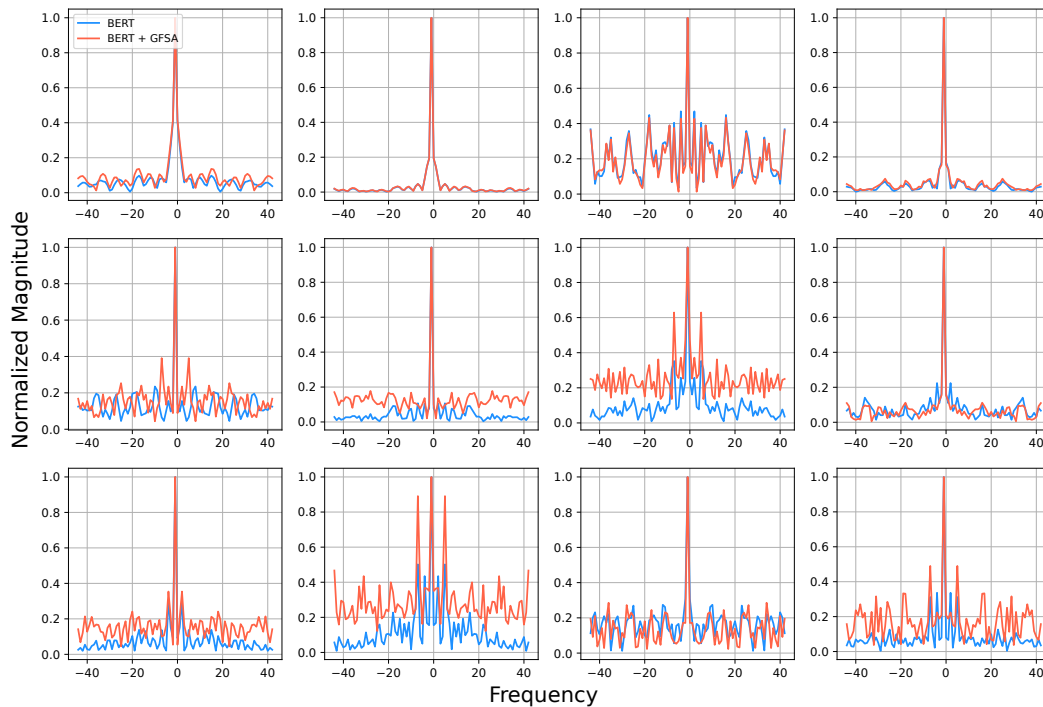


Figure 7: Visualization of the frequency responses for all 12 layers of BERT trained on STS-B dataset. The top-left figure corresponds to the first layer, and the bottom-right figure corresponds to the last layer.

E Frequency Analyses in the Singular Value Domain

Graph signal processing (GSP) [67, 68] can be understood as a generalized concept of DSP — in other words, DSP is a special case of GSP where a *line graph with n nodes* is used and therefore, the graph Fourier transform (GFT) of the line graph is identical to the discrete Fourier transform.

In the definition of GFT, we assume that the graph shift operator (GSO) $\mathbf{S}$ is diagonalizable. Considering the eigendecomposition of the GSO $\mathbf{S} = \mathbf{V}^\top \mathbf{\Lambda} \mathbf{V}$ with eigenvector $\mathbf{V}$, we can write the graph filter output as follows:

$$\mathbf{y} = \sum_{k=0}^K w_k \mathbf{S}^k \mathbf{x} = \sum_{k=0}^K \mathbf{V}^\top w_k \mathbf{\Lambda}^k \mathbf{V} \mathbf{x} = \mathbf{V}^\top \left(\sum_{k=0}^K w_k \mathbf{\Lambda}^k \right) \mathbf{V} \mathbf{x} = \mathbf{V}^\top g(\mathbf{\Lambda}) \mathbf{V} \mathbf{x}, \quad (11)$$

where $\mathbf{x} \in \mathbb{R}^n$ is a 1-dimensional graph signal, $\mathbf{\Lambda}$ is a diagonal matrix with eigenvalues, and $w_k \in [-\infty, \infty]$ is a coefficient.

However, one can use the singular value decomposition, when the GSO is not diagonalizable but symmetrically normalized, instead of the eigendecomposition [49]. Both the singular value decomposition and the eigendecomposition project the original signal onto a set of basis, but they use different basis sets. In the singular value decomposition, we sort the set of basis in ascending order of their eigenvalues, and perform frequency domain-like analyses [106, 49].

Since the self-attention matrix's row-wise sum is always 1, the following is the case: $\bar{\mathbf{A}} = \mathbf{D}^{-1} \mathbf{A} = \frac{1}{n} \mathbf{A}$, where n is the number of tokens. Maskey et al. [49] define the following symmetrically normalized adjacency matrix (SNA): $\mathbf{D}_{in}^{-1/2} \mathbf{A} \mathbf{D}_{out}^{-1/2}$. Since the degree of every node is n in the self-attention matrix, the following is the case: $\mathbf{D}_{in}^{-1/2} \mathbf{A} \mathbf{D}_{out}^{-1/2} = \mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} = \frac{1}{\sqrt{n}} \mathbf{A} \frac{1}{\sqrt{n}} = \frac{1}{n} \mathbf{A} = \bar{\mathbf{A}}$. Therefore, the self-attention matrix is a special case of SNAs.

F Matrix Polynomial vs. Graph Fourier Transform

There are two paradigms of implementing graph filters: i) matrix polynomial, which does not require diagonalizability, and ii) graph Fourier transform, which uses the eigendecomposition for diagonalizable adjacency matrices or uses the Jordan decomposition or the singular value decomposition for non-diagonalizable adjacency matrices.

Those two paradigms have their own weaknesses: i) the matrix polynomial approach requires explicit matrix multiplications, and ii) the graph Fourier transform approach requires expansive spectral decompositions. The matrix polynomial is preferred when there are not many matrix multiplications. Otherwise, the graph Fourier transform approach may be better since the matrix multiplication can be simplified after the decomposition.

Among those two, we use the first matrix polynomial approach with only three non-zero coefficients $\{w_0, w_1, w_K\}$ since it does not require the complicated spectral decomposition. Since we do not rely on any explicit spectral decomposition but on the matrix polynomial, any adjacency matrix can be used.

G Proof of Theorem 3.1

Theorem 3.1 (Filter characteristics based on coefficient values). *Let $\bar{\mathbf{A}}$ be a self-attention matrix interpreted as a graph with connected components. Consider the polynomial graph filter defined by $\sum_{k=0}^K w_k \bar{\mathbf{A}}^k$, where $w_2, w_3, \dots, w_{K-1} = 0$ and only w_0, w_1 , and w_K are non-zero. If the coefficients w_k for $k = 0, 1, K$ are positive and their sum is 1, then the polynomial filter acts as a low-pass filter, attenuating high-frequency components and promoting smoothness across the graph. Conversely, if $w_k = (-\alpha)^k$ for $k = 0, 1, K$ and $\alpha \in (0, 1)$ with sufficient large K , the polynomial filter exhibits high-pass filter behavior.*

Note that without filtering, the singular value ratio is $|\sigma_i^0|/|\sigma_1^0| = 1$. In the case where $|g(\sigma_i)/g(\sigma_1)| < 1 \ \forall i \geq 2$, it implies that after applying the graph filter g , the lowest frequency component further dominates, indicating that the graph filter acts as a low-pass filter. Conversely, in

the case where $|g(\sigma_i)/g(\sigma_1)| > 1 \forall i \geq 2$, it implies that after applying the graph filter g , the lowest frequency component σ_i no longer dominates, indicating that the graph filter acts as a high-pass filter.

Proof. We prove the low-pass filter result. For the case where w_0, w_1 , and w_K are positive and their sum is 1, we show that

$$|g(\sigma_1)| = |w_0 + w_1 + w_K| = 1 \quad (12)$$

Hence, proving Theorem G is equivalent to show $|g(\sigma_i)| < 1$.

$$|g(\sigma_i)| = |w_0 + w_1\sigma_i + w_K(\sigma_i + (K-1)(\sigma_i^2 - \sigma_i))| < |w_0 + w_1\sigma_i + w_K\sigma_i| = \sigma_i < 1 \quad (13)$$

since $\sigma_i + (K-1)(\sigma_i^2 - \sigma_i) = \sigma_i((K-1)\sigma_i - (K-2)) < \sigma_i((K-1) - (K-2)) = \sigma_i$

For the high-pass filter result, when $w_k = (-\alpha)^k/(k+1)$ where $k = 0, 1, K$ and $\alpha \in (0, 1)$, then we show that when $w_0 = 1, w_1 = -\alpha/2$,

$$\left| \frac{\lim_{K \rightarrow \infty} g(\sigma_i)}{\lim_{K \rightarrow \infty} g(\sigma_1)} \right| = \left| \frac{\lim_{K \rightarrow \infty} w_0 + w_1\sigma_i + w_K(\sigma_i + (K-1)(\sigma_i^2 - \sigma_i))}{\lim_{K \rightarrow \infty} w_0 + w_1 + w_K(1 + (K-1)(1-1))} \right| \quad (14)$$

$$= \left| \frac{\lim_{K \rightarrow \infty} 1 - \frac{\alpha}{2}\sigma_i + \frac{(-\alpha)^K}{(K+1)}(\sigma_i + (K-1)(\sigma_i^2 - \sigma_i))}{\lim_{K \rightarrow \infty} 1 - \frac{\alpha}{2} + \frac{(-\alpha)^K}{(K+1)}} \right| \quad (15)$$

$$= \left| \frac{1 - \frac{\alpha}{2}\sigma_i + (\lim_{K \rightarrow \infty} \frac{(-\alpha)^K}{(K+1)}(\sigma_i + (K-1)(\sigma_i^2 - \sigma_i)))}{1 - \frac{\alpha}{2}} \right| \quad (16)$$

$$= \left| \frac{1 - \frac{\alpha}{2}\sigma_i}{1 - \frac{\alpha}{2}} \right| > 1 \quad (17)$$

since

$$\lim_{K \rightarrow \infty} \frac{(-\alpha)^K}{(K+1)}(\sigma_i + (K-1)(\sigma_i^2 - \sigma_i)) = \lim_{K \rightarrow \infty} \frac{(-\alpha)^K}{(K+1)}(K-1)(\sigma_i^2 - \sigma_i) \quad (18)$$

$$= \lim_{K \rightarrow \infty} (-\alpha)^K \frac{(K+1)}{(K-1)}(\sigma_i^2 - \sigma_i) = 0 \quad (19)$$

It shows that the graph filter with $w_k = (-\alpha)^k/(k+1)$ for $k = 1, 2, K$ emphasizes high-frequency components and acts as a high-pass filter.

This proof supports that the behavior of the polynomial filter as either a low-pass or high-pass filter directly depends on the sign and values of the coefficients, as specified in Theorem 3.1. $\square$

H Proof of Theorem 4.1

Proof. The Frobenius norm of the self-attention is directly related to how far the softmax probabilities are from being uniform. For any matrix $\mathbf{M} \in \mathbb{R}^{m \times n}$, we have

$$\|\text{softmax}(\mathbf{M})\|_F = \sqrt{\frac{m + \sum_{i=1}^m d_{\chi^2}(S_i, U_n)}{n}}, \quad (20)$$

where S_i is the i -th row of $\text{softmax}(\mathbf{M})$, U_n is the uniform distribution over n elements, and $d_{\chi^2}(p, q) = \sum_i q_i(p_i/q_i - 1)^2$ is the χ^2 -divergence between p and q [14]. The Frobenius norm is maximized when the whole mass of the probabilities is on one element, which is a case for $d_{\chi^2}(S_i, U_n) = n - 1$ and $\|\text{softmax}(\mathbf{M})\|_F = \sqrt{m}$. Therefore, we can calculate the upper bound of Frobenius norm for $\bar{\mathbf{A}}$ as follows:

$$\|\bar{\mathbf{A}}\|_F \leq \sqrt{n}. \quad (21)$$

Note that $\bar{\mathbf{A}} \in \mathbb{R}^{n \times n}$ is a right-stochastic matrix normalized with row-wise softmax: i) all the elements of $\bar{\mathbf{A}}$ lie within $[0, 1]$, and ii) the row-sum in $\bar{\mathbf{A}}$ is equal to 1. Since the self-attention matrix

is a right-stochastic matrix, the power of the self-attention is also a right-stochastic matrix. Therefore, Eq. (21) is also hold for $\bar{\mathbf{A}}^K$ as follows:

$$\|\bar{\mathbf{A}}^K\|_F = \sqrt{\sum_{i,j} \bar{\mathbf{A}}_{i,j}^K} \leq \sqrt{\sum_{i,j} \bar{\mathbf{A}}_{i,j}} = \|\bar{\mathbf{A}}\|_F \leq \sqrt{n}. \quad (22)$$

Now, considering the error term E_K as given by Theorem 4.1, and applying the triangle inequality for matrix norms:

$$E_K = \|\bar{\mathbf{A}}^K - (\bar{\mathbf{A}} + (K-1)(\bar{\mathbf{A}}^2 - \bar{\mathbf{A}}))\|_F \quad (23)$$

$$\leq \|\bar{\mathbf{A}}^K\|_F + \|\bar{\mathbf{A}}\|_F + (K-1)(\|\bar{\mathbf{A}}^2\|_F + \|\bar{\mathbf{A}}\|_F) \quad (24)$$

$$\leq \sqrt{n} + \sqrt{n} + (K-1)(\sqrt{n} + \sqrt{n}) = 2\sqrt{n}K. \quad (25)$$

□

I Implementation of GFSA

The pseudo code of our GFSA is shown in Algorithm 1. For implementation, w_0 and w_1 can be set as hyperparameters optionally.

Algorithm 1 PyTorch-style pseudocode for GFSA

```
w_0 = torch.zeros(h)
w_1 = torch.ones(h)
w_K = torch.zeros(h)
I = torch.eyes(n) [None, None, ...]

def GFSA (att, K)
    att: original self-attention
    att_K: high order term
    att_K = att + (K-1) * (torch.mm(att,att)-att)
    gf_att: GFSA attention
    gf_att = w_0[None, :, None, None] * I
            + w_1[None, :, None, None] * att
            + w_K[None, :, None, None] * att_K
    return gf_att
```

J Comparison with Actual and Approximated High-order Terms

To compare the impact of the actual $\bar{\mathbf{A}}^K$ and the approximated $\bar{\mathbf{A}}^K$ in terms of accuracy, we experimented with BERT on GLUE and the results are summarized in Table 9. BERT_{BASE}+ $\bar{\mathbf{A}}^K$ denotes using the exactly calculated $\bar{\mathbf{A}}^K$ instead of the approximated $\bar{\mathbf{A}}^K$.

Table 9: Comparison of performance using the exactly calculated $\bar{\mathbf{A}}^K$ vs. the approximated $\bar{\mathbf{A}}^K$ for GLUE tasks

Datasets	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	56.79	93.81	88.70	88.32	88.16	84.96/84.15	91.63	66.06	82.51
+ GFSA (approximated $\bar{\mathbf{A}}^K$)	110M	59.56	94.15	90.60	88.46	88.33	85.12/85.06	91.95	68.95	83.58
+ GFSA (actual $\bar{\mathbf{A}}^K$)	110M	59.85	94.27	89.80	88.43	88.32	84.95/84.89	91.76	68.23	83.39

K Natural Language Understanding

K.1 Detailed Experimental Settings

We integrate GFSA into 3 pre-trained large language models: BERT, ALBERT, and RoBERTa. We evaluate them on the GLUE benchmark, which includes 3 categories of natural language understanding tasks: i) single-sentence tasks CoLA and SST-2; ii) similarity and paraphrasing tasks MRPC, QQP, and STS-B; iii) natural language inference tasks MNLI, QNLI, and RTE. For MNLI task, we experiment on both the matched (MNLI-m) and mismatched (MNLI-mm) versions. Following Devlin et al. [16], we report Matthews correlation for CoLA, F1 scores for QQP and MRPC, Spearman correlations for STS-B, and accuracy scores for the other tasks. For each task, we select the best hyperparameters for GFSA, and the other hyperparameters are fixed. We compare our GFSA with ContraNorm [28], one of the related methods that address oversmoothing. We finetune ContraNorm with the recommended hyperparameters in Guo et al. [28]. We initialize with a pre-trained language model and finetune with added GFSA for 5 epochs.

Dataset. The benchmark datasets we used are listed below.

- **CoLA.** The Corpus of Linguistic Acceptability [85] consists of English acceptability judgments drawn from books and journal articles. The target task is a binary classification task, and each sentence is determined to be grammatically acceptable or not.
- **SST-2.** The Stanford Sentiment Treebank [73] is a dataset in which each sentence is sourced from movie reviews and accompanied by human annotations of their sentiment. The target task is to classify binary sentiments for a single sentence.
- **MRPC.** The Microsoft Research Paraphrase Corpus [17] is a corpus of sentence pairs, which are automatically extracted from online news sources and annotated by humans. The target is to determine whether the sentences in the pair are semantically equivalent.
- **QQP.** The Quora Question Pairs [11] dataset is a collection of question pairs from the community question-answering website Quora. The target is to determine whether the questions in the pair are semantically equivalent.
- **STS-B.** The Semantic Textual Similarity Benchmark [8] is a collection of sentence pairs drawn from news headlines, video and image captions, and natural language inference data with human annotation. The target is a regression task to predict a similarity score from 0 to 5.
- **MNLI.** The Multi-Genre Natural Language Inference Corpus [87] is a crowdsourced collection of sentence pairs with textual entailment annotations. Given a premise sentence and a hypothesis sentence, the task is to predict whether the premise entails the hypothesis (entailment), contradicts the hypothesis (contradiction), or neither (neutral). The standard test set consists of private labels from the authors and evaluates both the matched (in-domain) and mismatched (cross-domain) sections.
- **QNLI.** The Stanford Question Answering [82] dataset is a question-answering dataset consisting of question-paragraph pairs, where one of the sentences in the paragraph contains the answer to the corresponding question written by an annotator. The task is to determine whether the context sentence contains the answer to the question.
- **RTE.** The Recognizing Textual Entailment [5] dataset comes from a series of annual textual entailment challenges. The target task is a binary entailment classification task.

BERT. BERT [16] consists with 12 layers, 12 heads, 768 hidden size, 512 maximum sequence length, and MLP dimension of 3072.

ALBERT. ALBERT [40] consists of 12 layers, 12 heads, 768 hidden dimensions, 512 maximum sequence length, 128 embedding dimensions, and MLP dimension of 3072.

RoBERTa. RoBERTa [44] consists of 12 layers, 12 heads, 768 hidden size, 514 maximum sequence length, and MLP dimension of 3072.

Training. For implementation, we adopt HuggingFace framework. We trained all models with 5 epochs with 32 batch size. The linear learning rate decay is used and initial learning rate is set to 2×10^{-5} . We use AdamW [46] optimizer, and weight decay is set to 0. All models are trained on 1 GPU and of NVIDIA RTX A5000 24GB.

K.2 Sensitivity to K

In this section, we explore the influence of the polynomial order, denoted as K , in our GFSA, conducting experiments on BERT_{BASE} finetuned with GLUE tasks. We search for values of K from 2 to 10, and the results are presented in Table 10. For each dataset, there is an optimal K and the performance of models using GFSA is generally robust to changes in K .

Table 10: Sensitivity results on various K with BERT_{BASE} finetuned on GLUE tasks

K	CoLA	SST2	MRPC	QQP	STSB	MNLI-m	MNLI-mm	QNLI	RTE
2	57.83	94.15	90.60	88.41	88.27	84.96	84.90	91.74	68.23
3	58.56	93.46	89.77	88.41	88.33	85.08	84.75	91.78	68.59
4	59.56	93.46	89.77	88.45	88.29	85.06	85.06	91.76	68.59
5	58.10	93.58	90.07	88.39	88.29	84.87	84.99	91.85	68.95
6	59.40	93.58	90.40	88.29	88.27	84.93	84.97	91.69	68.23
7	59.12	94.04	90.48	88.43	88.26	85.12	84.94	91.82	68.94
8	58.58	93.69	90.12	88.46	88.24	84.92	84.81	91.95	68.95
9	58.88	93.46	89.54	88.41	88.26	85.06	91.67	67.87	85.04
10	59.31	93.35	89.98	88.41	88.30	84.84	91.73	68.59	84.93

L Causal Language Modeling

L.1 Detailed Experimental Settings

Dataset. The benchmark datasets we used are listed below.

- **PTB.** Penn Treebank [47] dataset is a collection of text documents that have been extensively annotated with linguistic information, primarily syntactic and grammatical structures.
- **WikiText.** WikiText [50] dataset is a collection of over 100 million tokens extracted from the set of verified good and featured articles on Wikipedia. Compared to the preprocessed version of Penn Treebank (PTB), WikiText-2 is over 2 times larger and WikiText-103 is over 110 times larger.

GPT2. GPT2 [61] is a Transformer pretrained on a very large corpus of English data in a self-supervised fashion without any human labelling on dataset. It automatically generate inputs and labels from those texts, and trained to guess the next word in sentences. For implementation, we adopt HuggingFace Framework ⁴. For all experiments, GPT2 has 12 layers with 12 attention heads, 768 hidden size and 1024 maximum sequence length, resulting in a total of 117 million parameters.

Training. We finetune GPT2 with 4 batch size, 5×10^{-5} learning rate and linear learning weight decay using adamW [46] optimizer. We also apply dropout with probability 0.1. Following [93], we train models for 15 epochs with PTB, 4 epochs with WikiText-103 and 10 epochs with WikiText-2. We use sensitivity metric, i.e., perplexity, which is a commonly used metric to evaluate the performance of language models, particularly in language modeling and text generation tasks. perplexity measures how well a language modeling can predict a sequence of words in a given text or a test dataset. All the experiments are conducted on 1 GPU and of NVIDIA RTX 3090 24GB.

L.2 Sensitivity to K

We conducted a sensitivity study on K of GPT-2 across all datasets, and the results are presented in Table 11. For PTB and WikiText-2, GFSA exhibits the best performance when K is high, typically

⁴<https://github.com/huggingface/transformers>

around 8 or 9. However, for WikiText-103, GFSA achieves the best perplexity when K is small, specifically when K is 3 or 4.

Table 11: Results comparison on GPT-2 finetuned with GFSA

Method	#Params	PTB	WikiText-2	WikiText-103
GPT2 [61]	117M	19.513	20.966	15.939
GPT2 + GFSA($K = 2$)	117M	19.459	20.929	15.920
GPT2 + GFSA($K = 3$)	117M	19.456	20.927	15.919
GPT2 + GFSA($K = 4$)	117M	19.453	20.927	15.919
GPT2 + GFSA($K = 5$)	117M	19.452	20.925	15.920
GPT2 + GFSA($K = 6$)	117M	19.451	20.925	15.920
GPT2 + GFSA($K = 7$)	117M	19.450	20.925	15.921
GPT2 + GFSA($K = 8$)	117M	19.450	20.924	15.921
GPT2 + GFSA($K = 9$)	117M	19.450	20.923	15.921

M Image Classification

M.1 Detailed Experimental Settings

Our code is implemented based on the timm library [86]. In the case of our training recipe, it is the same as experimental setting of Wang et al. [80] that follows the training recipes of Touvron et al. [74] and Touvron et al. [75]. To apply our GFSA to existing base models such as DeiT, CaiT, and Swin, we consider a range of K between 2 and 5. For 12-layer DeiT, we follow the same hyperparameters from Wang et al. [80]. We set the dropout rate to 0 and 0.2 for 12-layer and 24-layer DeiT, respectively. For CaiT, we apply our GFSA on only to the patch embedding layer. All other hyper-parameters are kept consistent with the original papers of DeiT [74], CaiT [75] and, Swin [45]. All models are trained on NVIDIA RTX 3090 24GB.

M.2 FLOPs & Throughput

In Table 12, we report the number of FLOPs and throughput. With GFSA plugged in, the FLOP count is either the same or no different. For DeiT-S with 24 layers, which shows a slight FLOP increase with GFSA plugged in. However, for the rest of the settings, the models have the same number of Flops. For throughput, it tends to decrease because calculating the high-order term is an additional cost.

Table 12: Experimental evaluation of GFSA plugged into DeiT-S, CaiT-S, and Swin-S

Backbone	Method	Input Size	#Layers	#Params	#FLOPs	#Throughput	Top-1 Acc
DeiT	DeiT-S	224	12	22.0M	4.57G	856.07	79.8
	DeiT-S + GFSA	224	12	22.0M	4.57G	614.54	81.1 ($\uparrow 1.3$)
	DeiT-S	224	24	43.3M	9.09G	423.68	80.5
	DeiT-S + GFSA	224	24	43.3M	9.10G	314.75	81.5 ($\uparrow 1.0$)
CaiT	CaiT-S	224	24	46.9M	9.34G	574.66	82.6
	CaiT-S + GFSA	224	24	47.0M	9.34G	406.96	82.8 ($\uparrow 0.2$)
Swin	Swin-S	224	24	49.6M	8.75G	912.38	82.9
	Swin-S + GFSA	224	24	49.6M	8.75G	714.60	83.0 ($\uparrow 0.1$)

M.3 Sensitivity to K

We also perform the sensitivity analysis for K . Tables 13 and 14 show the results of sensitivity analysis for DeiT-S and CaiT-S with GFSA plugged in. For 12-layer DeiT-S, GFSA performance of 81.12 is highest when $K = 3$. When the GFSA has a K of 2, the performance is worse than the

original DeiT-S, but when the K is 3 or higher, the performance is better than the original DeiT-S, and most surprisingly, the performance is better than the 24-layer DeiT-S.

CaiT-S shows the highest performance of 82.84 when $K = 4$. For CaiT-S, the accuracy is slightly lower than that of the original CaiT-S when $K = 2$, but it starts to exceed the accuracy of CaiT-S when K is 3 or higher.

Table 13: Sensitivity to K for 12-layer DeiT-S + GFSA

K	2	3	4	5
Top-1 Acc (%)	79.27	81.12	80.86	81.07

Table 14: Sensitivity to K for 24-layer CaiT-S + GFSA

K	2	3	4
Top-1 Acc (%)	82.54	82.65	82.84

M.4 Full Experimental Results

In Table 15, we consider all three classes, CNN only, CNN + Transformer, and pure Transformer, to compare more different models than in Table 3. In particular, in the Transformer category, we only test with lightweight models with similar number of parameters, such as ViT-S and DeiT-S. Compared to existing techniques, the improvements by GFSA already surpasses LayerScale (0.7%) [75], LateInsertion (0.6%) [75], and HAT [4] (1.38%).

Table 15: Compared with state-of-the-art models on ImageNet-1k dataset. The number in ($\uparrow$) indicates the performance improvement over the base model.

Category	Method	Input Size	#Layers	#Params	Top-1 Acc
CNN	ResNet-152 [29]	224	152	230M	78.1
	DenseNet-201 [34]	224	201	77M	77.6
CNN + Transformer	CVT-21 [88]	224	21	32M	82.5
	Refiner [102]	224	16	86M	81.2
Transformer	ViT-S/16 [19]	224	12	49M	78.1
	ViT-B/16 [19]	224	12	86M	79.8
	DeiT-S [74]	224	12	22M	79.8
	DeiT-S + LayerScale [75]	224	12	22M	80.5
	DeiT-S + LateInsertion [75]	224	12	22M	80.5
	DeiT-S + ClassAttention [75]	224	12	22M	80.6
	DeiT-S + AttnScale [80]	224	12	22M	80.7
	DeiT-S + FeatScale [80]	224	12	22M	80.9
	DeiT-S + HAT [4]	224	12	22M	80.9
	DeiT-S + Diverse [10]	224	12	22M	80.6
	DeiT-S + ContraNorm [28]	224	12	22M	80.4
	Swin-S [45]	224	12	50M	82.9
	T2T-ViT-24 [96]	224	24	64M	82.3
	DeepViT-24B [101]	224	24	36M	80.1
	DeiT-S [74]	224	24	43M	80.5
	CaiT-S [75]	224	24	47M	82.6
	DeiT-S + DiversePatch [25]	224	24	44M	82.2
	DeiT-S + LayerScale [75]	224	24	44M	82.4
	DeiT-S + AttnScale [80]	224	24	44M	81.1
	DeiT-S + FeatScale [80]	224	24	44M	81.3
	DeiT-S + ContraNorm [28]	224	24	43M	80.7
GFSA	DeiT-S + GFSA	224	12	22M	81.1 ($\uparrow$ 1.3)
	DeiT-S + GFSA	224	24	43M	81.5 ($\uparrow$ 1.0)
	CaiT-S + GFSA	224	24	47M	82.8 ($\uparrow$ 0.2)
	Swin-S + GFSA	224	12	50M	83.0 ($\uparrow$ 0.1)

M.5 Additional Comparison with SOTA Models

Our main experiment aims to determine whether introducing the GFSA layer would help improve performance in a base model, such as DeiT. We also compare our method with the recent models: SpectFormer [58], SVT [57], NeuTRENO [52], FNet [42], and GFNet [64]. We use a 12-layer setup to ensure a fair comparison in Table 16.

GFNet [64] can reduce the number of parameters, but there is a performance penalty. However, the performance improvement of DeiT-S+GFSA is relatively greater than DeiT-S compared to other models. SpectFormer [58] and SVT [57] have advantages in calculation amount and model complexity, and performance is improved over DeiT-S, but Top-1 and Top-5 accuracies are lower than those using GFSA. Additionally, NeuTRENO [52] also improves as much as GFSA compared to DeiT-S, but GFSA still has higher Top-1 accuracy.

Table 16: Compared with state-of-the-art models on ImageNet-1k

Method	Input Size	#Layers	#Params	Top-1 Acc	Top-5 Acc
DeiT-S	224	12	22M	79.8	95.0
Fnet-XS [42]	224	12	20M	71.2	-
GFNet-XS [64]	224	12	16M	78.6	94.2
SpectFormer-XS [58]	224	12	20M	80.2	94.7
SVT-XS [57]	224	12	20M	79.9	94.5
DeiT-S + NeuTRENO [52]	224	12	20M	80.7	95.4
DeiT-S + GFSA	224	12	22M	81.1	95.4

M.6 Additional Experiments with Guo et al. [28]’s setting

To make a fair comparison with ContraNorm [28], one of the related studies that mitigates over-smoothing, we run additional experiments to match their experimental setup.

Setting. We follow the training recipe used by Guo et al. [28], which is a slightly modified version of Touvron et al. [74]’s recipe. Guo et al. [28] use AdamW optimizer with cosine learning rate decay. We select the DeiT-T and DeiT-S for ImageNet-1k. “T” and “S” denote tiny and small model sizes, respectively. For all experiments, the image size is set to be 224x224. We train each model for 300 epochs and the batch size is set to 1024. For ContraNorm, we train with their recommended hyperparameters. All models are trained on 4 GPUs and of NVIDIA RTX A6000 48GB.

Results. In Table 17, DeiT-T and DeiT-S with GFSA outperform vanilla DeiT-T and DeiT-S in all layer settings. GFSA improves the performance of DeiT-T with 12 layers by 1.52%. The largest gain is a 4.88% improvement on 16-layer DeiT-T. This shows that the effect of GFSA is larger than the effect of ContraNorm. For DeiT-S with 16 layers, surprisingly, GFSA is able to increase the performance by 80.83%, meaning that GFSA brings benefits with a 3.23% improvement.

Table 17: Experiment results on ImageNet-1k

Method	#Layers=12	#Layers=16	#Layers=24
DeiT-T	76.52	75.34	76.76
DeiT-T + ContraNorm	77.03	78.72	78.12
DeiT-T + GFSA	77.68	79.02	78.64
DeiT-S	77.32	78.25	77.69
DeiT-S + ContraNorm	77.80	79.04	78.67
DeiT-S + GFSA	79.86	80.83	79.15

Sensitivity to K . In Table 18, we experiment with a sensitivity analysis for K . For DeiT-T, the performance of GFSA generally improves when K is 4 or 5. On the other hand, GFSA performs better at lower K for settings that are layers 16 and 24 for DeiT-S.

Table 18: Varying K for DeiT-T and DeiT-S

Method	K	#Layers=12	#Layers=16	#Layers=24
DeiT-T + GFSA	2	76.92	78.14	78.40
DeiT-T + GFSA	3	77.41	77.76	78.09
DeiT-T + GFSA	4	77.01	79.02	78.64
DeiT-T + GFSA	5	77.68	78.14	78.64
DeiT-S + GFSA	2	79.84	80.83	79.15
DeiT-S + GFSA	3	79.85	79.39	79.07
DeiT-S + GFSA	4	79.86	79.44	79.10

N Automatic Speech Recognition

N.1 Detailed Experimental Settings

Dataset. We conduct experiments on the LibriSpeech ⁵ dataset [55], which consists of audio recordings paired with their transcriptions. The LibriSpeech dataset has approximately 1,000 hours of read English speech with a sampling rate of 16 kHz. We keep the original 16,000Hz sampling rate and compute 80-dim log-Mel filterbanks for a 25ms sliding window, strided by 10ms. The filterbank features are then normalized to zero mean and unit variance per input sequence. For implementation, we follow the recipes of SpeechBrain [65].

Evaluation Metric. Word error rate (WER (%)) is derived from the Levenshtein distance and compares a reference to a hypothesized word-level transcription. It is calculated by summing the number of word insertions, deletions, substitutions and dividing it by the total number of words in the reference transcription.

Vanilla Transformer. We use a vanilla Transformer to apply our GFSA. For implementation, we use a SpeechBrain [65] framework. The vanilla Transformer consists of i) 1D convolution to perform striding, ii) Transformer encoder with 12 layers, 4 heads, embedding dimension of 512, MLP dimension of 2048, and post-LayerNorm iii) decoder with 6 layers, 4 heads, embedding dimension of 512, MLP dimension of 2048, joint beamsearch, and iv) external Transformer language model with 12 layers, 12 heads, embedding dimension of 768, and MLP dimension of 3072.

Branchformer. We use one of the SOTA models, Branchformer [59] to plug-in our GFSA. Branchformer has two parallel branches, one for capturing global interactions using attention and the other for more localized context using convolutional gating MLP. The Branchformer architecture for speech recognition consists of i) 1D convolution to perform striding, ii) Branchformer encoder with 18 layers, 8 heads, embedding dimension of 512, and MLP dimension of 3072, iii) decoder with 6 layers, 8 heads, embedding dimension of 512, a convolutional spatial gating unit (CSGU) dimension of 3072, joint beamsearch, and iv) external Transformer language model with 12 layers, 12 heads, embedding dimension of 768, and MLP dimension of 3072.

Training. We follow a training recipe from SpeechBrain [65]. The standard LibriSpeech validation sets (dev-clean and dev-other) are used to tune all parameters and select the best models. Test sets (test-clean and test-other) are used only to report final WER performance. We train the pure Transformer for 100 epochs and the Branchformer for 120 epochs with a batch size of 16. We use a data augmentation method on all models using SpecAugment [56]. SpecAugment applies time and frequency masking as well as time warping to the input spectrum. For Branchformer, we use AdamW [46] optimizer with 0.9 and 0.98 coefficients for computing running averages of gradient and its square. The learning rate and weight decay in all models are 0.0008 and 0.01, respectively. We use a connectionist temporal classification (CTC) loss [26, 100]. We also apply dropout with probability 0.1 and label smoothing with weight 0.1 to mitigate overfitting. We fix the random seed as 74443 on all experiments. All models are trained on 1 GPU and of NVIDIA RTX A6000 48GB.

⁵<http://www.openslr.org/12>

Hyperparameters. In Table 19, we describe main hyperparameters used in the automatic speech recognition task. For Transformer+GFSA and Branchformer+GFSA, we also report the best K hyperparameter.

Table 19: Main hyperparameters used in ASR

Model	Experimental Setting
Transformer	Encoder: Transformer (12 layers) Decoder: Transformer (6 layers) + (CTC/ATT joint) beamsearch + TransformerLM Augmentation: SpecAugment Features: 40 fbanks Pretraining: no Dropout: 0.1 Batchnorm: yes Number of epochs: 100 Batch size: 32 Learning rate: 0.0008 LR scheduler: Noam Optimizer: Adam Loss: CTC + KLdiv (Label Smoothing loss) CTC weight: 0.3
Transformer+GFSA	Encoder: Transformer (12 layers) Decoder: Transformer (6 layers) + (CTC/ATT joint) beamsearch + TransformerLM Augmentation: SpecAugment Features: 40 fbanks Pretraining: no Dropout: 0.1 Batchnorm: yes Number of epochs: 100 Batch size: 32 Learning rate: 0.0008 LR scheduler: Noam Optimizer: Adam Loss: CTC + KLdiv (Label Smoothing loss) CTC weight: 0.3 K : 2
Branchformer	Encoder: Branchformer Decoder: Transformer (6 layers) + (CTC/ATT joint) beamsearch + TransformerLM Augmentation: SpecAugment Features: 40 fbanks Pretraining: no Dropout: 0.1 Batchnorm: yes Number of epochs: 120 Batch size: 16 Learning rate: 0.0008 LR scheduler: Noam Optimizer: AdamW with coefficients 0.9 and 0.98 Loss: CTC + KLdiv (Label Smoothing loss) CTC weight: 0.3
Branchformer+GFSA	Encoder: Branchformer Decoder: Transformer (6 layers) + (CTC/ATT joint) beamsearch + TransformerLM Augmentation: SpecAugment Features: 40 fbanks Pretraining: no Dropout: 0.1 Batchnorm: yes Number of epochs: 120 Batch size: 16 Learning rate: 0.0008 LR scheduler: Noam Optimizer: AdamW with coefficients 0.9 and 0.98 Loss: CTC + KLdiv (Label Smoothing loss) CTC weight: 0.3 K : 3

N.2 Training Curve

We compare the training and validation curves for LibriSpeech 100h in Fig. 8. The training loss curve of GFSA is lower than the pure Transformer. GFSA stabilizes the loss curve of pure Transformer slightly earlier.

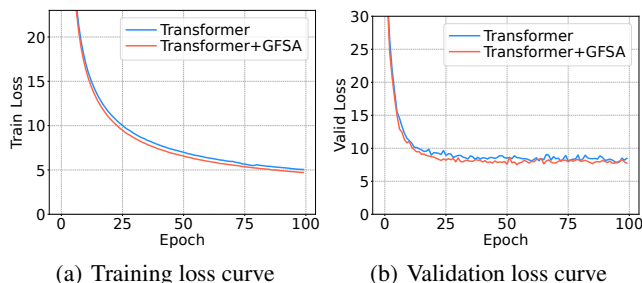


Figure 8: Training curve on LibriSpeech 100h

O Graph-level Tasks

O.1 Detailed Experimental Settings

Experimental settings for Graphormer. We describe benchmark datasets and Graphormer as the backbone model we used. ZINC [35] is the most popular real-world molecular dataset to predict graph property regression for constrained solubility, an important chemical property for designing generative GNNs for molecules. Uniform sampling is adopted for data splitting. We use a ZINC-subset of small-scale dataset. PCQM4M-LSC [33] is 2D molecular graphs, which is one of the most practically relevant quantum chemical properties of molecule science. The task is to predict density functional theory (DFT)-calculated HOMO-LUMO energy gap of molecules given their graphs. PCQM4M-LSC is unprecedentedly large in scale comparing to other labeled graph-level prediction datasets, which contain more than 3.8M graphs. We use PCQM4M and PCQM4Mv2 large-scale datasets.

Following Graphormer [94], we use Graphormer for PCQM4M and Graphormer_{SLIM} for ZINC. Graphormer consists of 12 encoder layers, 80 encoder embedding dimension, and 768 MLP dimension. It employs 32 encoder heads and 24 hidden dimension for each head. Graphormer_{SLIM} consists of 12 encoder layers, 80 encoder embedding dimension, and 80 MLP dimension. It employs 8 encoder heads and 10 hidden dimension for each head. We use adamW [46] optimizer with 0.9 and 0.999 coefficients for running averages of gradient and its square, and use Mean Absolute Error (MAE) as loss function. We use polynomial learning rate decay, with initial learning rate set to 2×10^{-4} and end learning rate set to 1×10^{-9} . For ZINC, we set batch size as 256, max epochs as 10k, and warm-up stage step as 40k. For PCQM4M and PCQM4Mv2, we set batch size as 1024, max epochs as 300, and warm-up stage step as 60k. All models are trained on 1 GPU and of NVIDIA RTX 3090 24GB. We conduct experiments with 4 different seeds.

Experimental settings for GPS and Graph-ViT. We use various benchmark datasets to experiment with our GFSA on GPS [63] and Graph-ViT [30]: Peptide-func and Peptide-struct from Long-Range Graph Benchmark (LRGB) [21], MNIST, CIFAR10, and ZINC from Benchmarking GNNs [22], and Moltox21 and Molhiv from OGB [32]. We fix all hyperparameters of GPS and Graph-ViT as recommended in their paper to ensure a fair comparison. To plug GFSA into GPS, we replace the self-attention module of GPS with GFSA. For Graph-ViT, we apply GFSA instead of the Hadamard self-attention mechanism and compare it with this self-attention method. We conduct experiments on GPS and Graph-ViT in their open code frameworks for a fair comparison:

- GPS: <https://github.com/rampasek/GraphGPS>
- Graph-ViT: <https://github.com/XiaoxinHe/Graph-ViT-MLPMixer>

O.2 Experimental Results with Standard Deviation

We conduct experiments following the experimental environments of Graphormer [94] using 4 different seeds. Due to space constraints, only the mean values are reported in Tables 4 and 5. In Tables 20 and 21, we report the results with mean and standard deviations.

Table 20: Experimental results and number of parameters on ZINC

Method	#Params	MAE
Graphormer	500K	0.1240 $\pm$ 0.006
Graphormer + GFSA	500K	0.1189 $\pm$ 0.002

Table 21: Experimental results and number of parameters on PCQM4M and PCQM4Mv2

Method	#Params	PCQM4M		PCQM4Mv2	
		Train	Validate	Train	Validate
Graphormer	48.3M	0.0535 $\pm$ 0.038	0.1286 $\pm$ 0.016	0.0250 $\pm$ 0.000	0.0862 $\pm$ 0.000
Graphormer + GFSA	48.3M	0.0312 $\pm$ 0.001	0.1193 $\pm$ 0.000	0.0249 $\pm$ 0.000	0.0860 $\pm$ 0.000

P Code Defect Detection

P.1 Detailed Experimental Settings

Dataset. We use Devign dataset provided by [105], which is a binary classification task to evaluate whether a C language code is vulnerable to software systems or not.

Implementation. We build our experiments on top of the open-sourced code ⁶ and recipes provided by Wang et al. [84].

RoBERTa. RoBERTa [44] is an encoder-only model trained with masked language modeling on code. All hyperparameters are consistent with the training method in the source code of Wang et al. [84].

PLBART. PLBART [2] is an encoder-decoder model based on BART [43] architecture. PLBART can support understanding and generation tasks. All hyperparameters are consistent with the training method in the source code of Wang et al. [84].

CodeBERT. CodeBERT [23] is a model trained on masked language modeling and replaced token detection. CodeBERT is a bimodal pretrained model based on Transformer with 12 layers for programming language and natural language. All hyperparameters are consistent with the training method in the source code of Wang et al. [84].

CodeT5. CodeT5 is an encoder-decoder framework with the same architecture as T5 [62]. It aims to derive generic representations for programming language and natural language via pre-training on unlabeled source code. CodeT5-small has 6 encoder layers, 6 decoder layers, 8 attention heads, 512 dimensional hidden states, and 60M parameters. The other models have 12 encoder layers, 12 decoder layers, 12 attention heads, 768 dimensional hidden states, and 220M parameters. All hyperparameters are consistent with the training method in the source code of Wang et al. [84].

Training. The pre-trained models mentioned above are applied to this downstream task. We add GFSA directly on top of self-attention. We finetune baselines and GFSA models for 10 epochs with a batch size of 16. We use early stopping strategy with a patience of 2. Models generate binary labels from unigram sequences at the decoder for defect detection task. We employ accuracy for evaluating the code defect detection task. All models are trained on 1 GPU and of NVIDIA RTX A6000 48GB.

⁶<https://github.com/salesforce/CodeT5>

P.2 Case Study

In Listing 1, we show one of case for code snippets of defects in QEMU⁷ that CodeT5-base does not predict correctly, but that *only* CodeT5-base+GFSA predicts. The commit message⁸ for this case is as follow:

Needed for changing cpu_has_work() argument type to CPUState, used in h_cede().

h_cede() is the hypercall that asks the hypervisor to shut down the CPU. Previously, this hypercall simply passed the CPUID, so the hypervisor did not know what state the CPU was in. This change allows the hypervisor to know whether the CPU is actually performing work. If the CPU is performing a task, the hypervisor waits for the CPU to complete the task.

In this context, accurately predicting defects like the one above is very important, and applying GFSA to CodeT5-base helps in terms of performance improvement.

```
1@@ -204,7 +204,7 @@ static target_ulong put_tce_emu(sPAPRTCETable *
    tcet, target_ulong ioba,
2- static target_ulong h_put_tce(CPUPPCState *env, sPAPREnvironment *
    spapr
3+ static target_ulong h_put_tce(PowerPCCPU *cpu, sPAPREnvironment *
    spapr
4
    , target_ulong opcode, target_ulong *
    args)
5 {
6     target_ulong liobn = args[0];
7     target_ulong ioba = args[1];
8     target_ulong tce = args[2];
9     VIOsPAPRDevice *dev = spapr_vio_find_by_reg(spapr->vio_bus, liobn);
10    VIOsPAPR_RTCE *rtce;
11    if (!dev) {
12        hcall_dprintf("LIOBN 0x" TARGET_FMT_lx " does not exist\n", liobn)
        ;
13        return H_PARAMETER;
14    }
15    ioba &= ~(SPAPR_VIO_TCE_PAGE_SIZE - 1);
16    #ifdef DEBUG_TCE
17        fprintf(stderr, "spapr_vio_put_tce on %s ioba 0x" TARGET_FMT_lx "
        TCE 0x" TARGET_FMT_lx "\n", dev->qdev.id, ioba, tce);
18    #endif
19    if (ioba >= dev->rtce_window_size) {
20        hcall_dprintf("Out-of-bounds IOBA 0x" TARGET_FMT_lx "\n", ioba);
21        return H_PARAMETER;
22    }
23    rtce = dev->rtce_table + (ioba >> SPAPR_VIO_TCE_PAGE_SHIFT);
24    rtce->tce = tce;
25    return H_SUCCESS;
26 }
```

Listing 1: An example commit history for defects in Devign dataset

Q Code Clone Detection

Q.1 Detailed Experimental Settings

Dataset. Code clone detection aims to measure the similarity between two code snippets and predict whether they have the same functionality. We experiment with the Java data provided by Wang et al. [83].

⁷<https://www.qemu.org>

⁸<https://github.com/qemu/qemu/commit/b13ce26d3e8c6682044ae84920f2417b30ce356b>

Implementation. We build our experiments on top of the open-sourced code⁹ and recipes provided by Wang et al. [84].

Training. We finetune both CodeT5 and CodeT5+GFSA for one epoch with a batch size of 16. We also use early stopping with patience of 2. CodeT5 and CodeT5+GFSA encode source code and take the representation to calculate similarity of two code snippets. We employ F1 score for evaluating this task. All models are trained on 1 GPU and of NVIDIA RTX A6000 48GB.

Q.2 Experiment Result

Table 22 shows results comparing CodeT5 and CodeT5 with GFSA. The result shows that by using our GFSA, CodeT5 models improve their performance. CodeT5-small+GFSA provides a 0.61% improvement over Code5T-small.

Table 22: Results on the code clone detection task

Method	Clone F1
CodeT5-small [84]	94.36
CodeT5-small + GFSA	94.94 ($\uparrow 0.61\%$)
CodeT5-base [84]	94.31
CodeT5-base + GFSA	94.92 ($\uparrow 0.64\%$)

R Time Complexity and Empirical Runtime Analysis

Time Complexity. The time complexity of original self-attention is $\mathcal{O}(n^2d)$. But our GFSA has a high order term. Therefore, the time complexity of GFSA has $\mathcal{O}(n^2d + n^3)$. If n is smaller than d , the time complexity approaches $\mathcal{O}(n^2d)$, which is the complexity of original self-attention.

Empirical Runtime Analysis. We report the training time of various methods with GFSA in Tables 23 to 28. In general, the training time of methods with GFSA is slightly longer than that of existing methods. For example, the Transformer for the automatic speech recognition task increases from 190.5 seconds to 191.6 seconds on Librispeech 100h dataset, as increases of only 1 second. Instead of computing higher-order polynomial terms, our GFSA approximates them, with only a small increase in runtime, which is not very significant.

Table 23: Training time (seconds per epoch) on GLUE tasks. *s* denotes the abbreviation for second. **Avg** denotes the average training time across all tasks.

Datasets	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	17s	182s	17s	1483s	24s	2004s	580s	18s	541s
BERT _{BASE} + GFSA	110M	19s	192s	19s	1571s	25s	2147s	621s	20s	577s
ALBERT _{BASE} [40]	11M	15s	188s	20s	1506s	25s	2072s	612s	19s	557s
ALBERT _{BASE} + GFSA	11M	16s	197s	21s	1604s	26s	2219s	659s	21s	595s
RoBERTa _{BASE} [44]	125M	17s	190s	18s	1492s	25s	2012s	593s	18s	546s
RoBERTa _{BASE} + GFSA	125M	19s	200s	19s	1580s	26s	2151s	634s	20s	581s

Table 24: Training time (seconds per epoch) on causal language modeling tasks.

Method	#Params	PTB	WikiText-2	WikiText-103	Avg
GPT2 [61]	117M	89.1s	195.7s	9638.4s	3307.8s
GPT2 + GFSA	117M	160.3s	354.2s	17424.6s	5979.7s

⁹<https://github.com/salesforce/CodeT5>

Table 25: Training time (seconds per epoch) on ImageNet-1k

Backbone	Method	#Layers	#Params	#FLOPs	#Throughput	Runtime
DeiT	DeiT-S	12	22.0M	4.57G	856.07	551s
	DeiT-S + GFSA	12	22.0M	4.57G	614.54	814s
	DeiT-S	24	43.3M	9.09G	423.68	1508s
	DeiT-S + GFSA	24	43.3M	9.10G	314.75	1798s
CaiT	CaiT-S	24	46.9M	9.34G	574.66	1530s
	CaiT-S + GFSA	24	47.0M	9.34G	406.96	1624s
Swin	Swin-S	24	49.6M	8.75G	912.38	1897s
	Swin-S + GFSA	24	49.6M	8.75G	714.60	1970s

Table 26: Training time (seconds per epoch) on graph-level tasks

Method	ZINC	PCQM4M	PCQM4Mv2
Graphormer [94]	9s	740s	817s
Graphormer + GFSA	9s	896s	955s

Table 27: Training time (seconds per epoch) on LibriSpeech datasets

Method	LibriSpeech 100h	LibriSpeech 960h
Transformer	190.5s	3049.3s
Transformer + GFSA	191.6s	3398.4s
Branchformer [59]	248.5s	4990.1s
Branchformer + GFSA	254.3s	4999.3s

Table 28: Training time (seconds per epoch) on the code defect prediction task

Method	Runtime
RoBERTa [44]	543.96s
RoBERTa + GFSA	537.79s
CodeBERT [23]	555.28s
CodeBERT + GFSA	561.43s
PLBART [2]	467.80s
PLBART + GFSA	470.19s
CodeT5-small [84]	301.11s
CodeT5-small + GFSA	309.04s
CodeT5-base [84]	362.28s
CodeT5-base + GFSA	373.22s

S Inference Time Analysis

We report the inference time of various methods with GFSA in Tables 29 to 34.

Table 29: Inference time on GLUE tasks. *s* denotes the abbreviation for second. **Avg** denotes the average training time across all tasks.

Datasets	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	1.0 <i>s</i>	1.4 <i>s</i>	1.2 <i>s</i>	48.7 <i>s</i>	1.9 <i>s</i>	15.5 <i>s</i>	10.1 <i>s</i>	1.2 <i>s</i>	10.0 <i>s</i>
BERT _{BASE} + GFSA	110M	1.1 <i>s</i>	1.4 <i>s</i>	1.2 <i>s</i>	52.3 <i>s</i>	2.0 <i>s</i>	16.8 <i>s</i>	11.0 <i>s</i>	1.3 <i>s</i>	11.0 <i>s</i>
ALBERT _{BASE} [40]	11M	1.1 <i>s</i>	1.6 <i>s</i>	1.4 <i>s</i>	58.4 <i>s</i>	2.2 <i>s</i>	18.4 <i>s</i>	12.1 <i>s</i>	1.3 <i>s</i>	12.0 <i>s</i>
ALBERT _{BASE} + GFSA	11M	1.2 <i>s</i>	1.7 <i>s</i>	1.4 <i>s</i>	62.1 <i>s</i>	2.3 <i>s</i>	19.7 <i>s</i>	13.1 <i>s</i>	1.4 <i>s</i>	13.0 <i>s</i>
RoBERTa _{BASE} [44]	125M	1.0 <i>s</i>	1.4 <i>s</i>	1.1 <i>s</i>	47.0 <i>s</i>	1.9 <i>s</i>	15.0 <i>s</i>	9.9 <i>s</i>	1.2 <i>s</i>	10.0 <i>s</i>
RoBERTa _{BASE} + GFSA	125M	1.1 <i>s</i>	1.4 <i>s</i>	1.2 <i>s</i>	50.4 <i>s</i>	2.0 <i>s</i>	16.3 <i>s</i>	10.8 <i>s</i>	1.2 <i>s</i>	11.0 <i>s</i>

Table 30: Inference time on causal language modeling tasks.

Method	#Params	PTB	WikiText-2	WikiText-103	Avg
GPT2 [61]	117M	3.2 <i>s</i>	7.4 <i>s</i>	7.4 <i>s</i>	6.0 <i>s</i>
GPT2 + GFSA	117M	5.5 <i>s</i>	12.9 <i>s</i>	12.9 <i>s</i>	10.4 <i>s</i>

Table 31: Inference time on ImageNet-1k

Backbone	Method	#Layers	Inference Time
DeiT	DeiT-S	12	52 <i>s</i>
	DeiT-S + GFSA	12	53 <i>s</i>
	DeiT-S	24	68 <i>s</i>
	DeiT-S + GFSA	24	69 <i>s</i>
CaiT	CaiT-S	24	105 <i>s</i>
	CaiT-S + GFSA	24	107 <i>s</i>
Swin	Swin-S	24	17 <i>s</i>
	Swin-S + GFSA	24	17 <i>s</i>

Table 32: Inference time on graph-level tasks

Method	ZINC	PCQM4M	PCQM4Mv2
Graphormer [94]	8 <i>s</i>	99 <i>s</i>	31 <i>s</i>
Graphormer + GFSA	8 <i>s</i>	117 <i>s</i>	29 <i>s</i>

Table 33: Inference time on LibriSpeech datasets

Method	LibriSpeech 100h	LibriSpeech 960h
Transformer	328.1 <i>s</i>	323.7 <i>s</i>
Transformer + GFSA	329.5 <i>s</i>	343.3 <i>s</i>
Branchformer [59]	299.4 <i>s</i>	328.7 <i>s</i>
Branchformer + GFSA	305.5 <i>s</i>	354.1 <i>s</i>

Table 34: Inference time on the code defect prediction task

Method	Inference Time
RoBERTa [44]	22.4 <i>s</i>
RoBERTa + GFSA	23.9 <i>s</i>
CodeBERT [23]	23.8 <i>s</i>
CodeBERT + GFSA	24.1 <i>s</i>
PLBART [2]	37.7 <i>s</i>
PLBART + GFSA	39.3 <i>s</i>
CodeT5-small [84]	78.2 <i>s</i>
CodeT5-small + GFSA	82.5 <i>s</i>
CodeT5-base [84]	83.2 <i>s</i>
CodeT5-base + GFSA	88.5 <i>s</i>

T Results of the Strategy for Efficiency

In Tables 35 to 39, the results show that this strategy can reduce the increase in runtime and maintain performance compared to using GFSA for all layers.

Table 35: Comparison of performance using GFSA on all layers vs. GFSA_{even} on even layers for GLUE tasks

Datasets	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	56.79	93.81	88.70	88.32	88.16	84.96/84.15	91.63	66.06	82.51
BERT _{BASE} + GFSA	110M	59.56	94.15	90.60	88.46	88.33	85.12/85.06	91.95	68.95	83.58
BERT _{BASE} + GFSA _{even}	110M	58.80	93.69	90.50	88.47	88.27	85.13/85.02	91.65	70.76	83.59

Table 36: Comparison of training time (seconds per epoch) using GFSA on all layers vs. GFSA_{even} on even layers for GLUE tasks

Datasets	#Params	CoLA	SST-2	MRPC	QQP	STS-B	MNLI-m/mm	QNLI	RTE	Avg
BERT _{BASE} [16]	110M	17s	182s	17s	1483s	24s	2004s	580s	18s	541s
BERT _{BASE} + GFSA	110M	19s	192s	19s	1571s	25s	2147s	621s	20s	577s
BERT _{BASE} + GFSA _{even}	110M	17s	185s	18s	1506s	24s	2061s	595s	18s	553s

Table 37: Comparison of performance using GFSA on all layers vs. GFSA_{even} on even layers for causal language modeling tasks

Method	#Params	PTB	WikiText-2	WikiText-103	Avg
GPT2 [61]	117M	19.513	20.966	15.939	18.806
GPT2 + GFSA	117M	19.450	20.923	15.919	18.764
GPT2 + GFSA _{even}	117M	19.453	20.926	15.923	18.767

Table 38: Comparison of training time (seconds per epoch) using GFSA on all layers vs. GFSA_{even} on even layers for causal language modeling tasks

Method	#Params	PTB	WikiText-2	WikiText-103	Avg
GPT2 [61]	117M	89.1s	195.7s	9638.4s	3307.8s
GPT2 + GFSA	117M	160.3s	354.2s	17424.6s	5979.7s
GPT2 + GFSA _{even}	117M	127.4s	279.1s	13761.4s	4722.6s

Table 39: Comparison of using GFSA on all layers vs. GFSA_{even} on even layers for ImageNet-1k

Method	Input Size	#Layers	#Params	Top-1 Acc	Runtime
DeiT-S	224	12	43M	79.8	551s
DeiT-S + GFSA	224	12	43M	81.1	814s
DeiT-S + GFSA _{even}	224	12	43M	81.0	595s

U GFSA in Linear Transformers

Table 40 shows the accuracy, runtime, and GPU usage results on Long Range Arena benchmark using ListOps and Image datasets.

Table 40: Comparison of accuracy (%), runtime (*s* per 1,000 steps) and GPU usage (GB) on Long Range Arena benchmark

Method	ListOps (2K)			Image (4K)		
	Accuracy	Runtime	GPU usage	Accuracy	Runtime	GPU usage
Transformer [77]	37.1	198.3	5.50	38.2	345.1	5.88
Transformer+GFSA	37.6	635.8	10.87	40.2	737.2	11.20
Linformer [81]	37.3	63.4	1.73	37.8	158.5	3.45
YOSO-E [97]	37.3	85.7	0.37	39.8	114.2	1.42
Efficient Attention [70]	36.9	49.2	0.57	40.2	121.1	1.14
Efficient Attention + GFSA	37.9	53.8	0.67	40.4	135.8	1.33

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We clearly stated our claims in the introduction, including the contributions of the paper.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: In Section 6, we stated the computational overhead as a limitation of our method and provided strategies to address it.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: We have included proofs for Theorem 4.1 and Theorem 3.1 in Appendix H and Appendix G.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Our reproducibility statements are reported in A.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: See the statement about reproducibility in Appendix A.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Experimental environments, datasets, training methods are reported in Section 5.1 to 5.6 and in Appendix K to 5.6.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We tried to follow the experimental setup for benchmarking each task. Among the tasks we considered, it is conventional to report results from four experiments in the graph-level tasks task, so we report the mean and standard deviation together in the Appendix O.2.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: Our paper specifies the type of computer resource information (e.g. GPU) for the experiment. We also indicated the amount of calculation or execution time for the experiment.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [\[Yes\]](#)

Justification: We have reviewed the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [\[Yes\]](#)

Justification: We discuss the potential positive and negative social impacts of our research in Appendix B.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: There is no such risk in our paper.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: In our paper, the authors or original owners of code, data, and models used are properly credited, and their licenses and terms of use are explicitly mentioned.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New Assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: The code we provided is well documented in the README and is anonymized.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and Research with Human Subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: Our research does not include crowdsourcing experiments and studies with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: Because our research does not involve human subjects, it is not subject to IRB approval.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.