
Enhancing Graph Classification Robustness with Singular Pooling

Sofiane Ennadir*

King AI Labs, Microsoft Gaming
sofiane.ennadir@king.com

Oleg Smirnov*

King AI Labs, Microsoft Gaming
oleg.smirnov@microsoft.com

Yassine Abbahaddou

LIX, École Polytechnique, IP Paris
yassine.abbahaddou@polytechnique.edu

Lele Cao

King AI Labs, Microsoft Gaming
lelecao@microsoft.com

Johannes F. Lutzeyer

LIX, École Polytechnique, IP Paris
johannes.lutzeyer@polytechnique.edu

Abstract

Graph Neural Networks (GNNs) have achieved strong performance across a range of graph representation learning tasks, yet their adversarial robustness in graph classification remains underexplored compared to node classification. While most existing defenses focus on the message-passing component, this work investigates the overlooked role of pooling operations in shaping robustness. We present a theoretical analysis of standard flat pooling methods (sum, average and max), deriving upper bounds on their adversarial risk and identifying their vulnerabilities under different attack scenarios and graph structures. Motivated by these insights, we propose *Robust Singular Pooling (RS-Pool)*, a novel pooling strategy that leverages the dominant singular vector of the node embedding matrix to construct a robust graph-level representation. We theoretically investigate the robustness of RS-Pool and interpret the resulting bound leading to improved understanding of our proposed pooling operator. While our analysis centers on Graph Convolutional Networks (GCNs), RS-Pool is model-agnostic and can be implemented efficiently via power iteration. Empirical results on real-world benchmarks show that RS-Pool provides better robustness than the considered pooling methods when subject to state-of-the-art adversarial attacks while maintaining competitive clean accuracy. Our code is publicly available at: <https://github.com/king/rs-pool>.

1 Introduction

Graph Neural Networks (GNNs) [24, 44, 38] have demonstrated strong performance on graph-structured data and are now a standard tool for learning node and graph representations. However, recent studies reveal that GNNs are susceptible to adversarial perturbations – small, deliberately designed changes to the adjacency matrix or node features that can drastically affect predictions [20]. This vulnerability raises concerns about their reliability and consequently, enhancing the adversarial robustness of GNNs has become a key research focus. Efforts include the design of novel attack methods [7, 49, 48, 17] and corresponding defense mechanisms [42, 46, 19]. However, prior work primarily investigates node-classification tasks, with limited attention paid to adversarial robustness

*Equal contribution.

in graph-classification settings [21]. Yet, robustness at the graph level is equally critical, particularly in domains where entire graphs, rather than individual nodes, constitute the prediction target. In drug discovery, for example, molecules are modeled as graphs, and predicting properties such as efficacy or toxicity is a graph classification problem [35, 22]. Adversarial perturbations, editing a bond type for instance, can lead to erroneous evaluations with serious implications in precision medicine. Similarly, in cybersecurity, classifying control-flow graphs of binary programs aids in malware detection [34]; subtle perturbations could allow malicious code to evade the anomaly classification.

Existing adversarial defenses for GNNs primarily focus on the message-passing component, leveraging techniques such as attention mechanisms [46], weight regularization [1, 10], and adjacency preprocessing [42]. While these approaches are applicable to both node and graph classification tasks, graph-level prediction introduces additional challenges that go beyond message passing [18], particularly in constructing global representations from structurally diverse inputs. This step relies on pooling mechanisms [5, 9], which aggregate node features into fixed-size graph embeddings. Pooling operations are generally categorized into two types [30]: *hierarchical pooling*, which performs multi-stage graph coarsening to produce progressively smaller subgraphs [5], and *flat pooling* (or readout functions), which compute a global representation in a single step using aggregation functions such as sum, average, or max [9]. While hierarchical pooling is a general architectural tool used throughout the network, flat pooling specifically produces the final representation used in graph-level classification. Given the importance of this final aggregation step, it is critical to investigate how flat pooling functions influence the robustness of GNNs.

This work offers a complementary perspective on GNN robustness by focusing on the pooling stage, which plays a critical role in bridging node-level representations and graph-level embeddings. We center our analysis on flat pooling methods, which are widely used for their simplicity, efficiency, and strong empirical performance. Through theoretical analysis, we examine how different pooling strategies influence the expected robustness and provide insights on selecting pooling functions that balance robustness and performance under varying attack scenarios.

To pursue this goal, we begin by formalizing adversarial attacks in the context of graph classification. This framework establishes an upper bound on the model’s expected adversarial risk, allowing us to quantify GNN vulnerability within defined neighborhoods of input graphs. We then theoretically analyze common pooling operations, derive their respective robustness bounds, and provide insights into their behavior under perturbations. These bounds reflect both graph structure and model-specific factors such as message-passing weights. Motivated by this analysis, we propose *Robust Singular Pooling (RS-Pool)*, a novel pooling mechanism that constructs graph-level representations using the dominant right-singular vector of the node embedding matrix. RS-Pool improves robustness by filtering out noise-sensitive components. RS-Pool is fully differentiable, model-agnostic, and compatible with existing GNN architectures and defense methods. Since, we rely solely on the dominant singular vector in our pooling operation, we do not require the expensive computation of the entire singular value decomposition and can instead employ the far more efficient power iteration method enabling practical deployment. Empirical results across multiple adversarial settings confirm that RS-Pool achieves strong robustness with minimal impact on clean accuracy and runtime performance.

The main contributions of this work are threefold. First, we introduce a theoretical framework that quantifies the effect of flat pooling operations on the resulting robustness of GNNs in graph classification. Second, we propose *RS-Pool*, a novel pooling mechanism that leverages the dominant singular vector of the node representation matrix to improve resilience to adversarial perturbations. Third, we develop an efficient approximation algorithm for RS-Pool and empirically validate its effectiveness in enhancing model robustness under adversarial attacks.

2 Related Work

The adversarial robustness of graph neural networks (GNNs) has attracted substantial attention [20, 37, 26, 11], spurring a variety of attack methodologies. These attacks are typically formulated as optimization problems that seek minimal perturbations of the adjacency matrix or node features to alter model predictions. They fall into two main categories: evasion attacks, which alter the graph after training [43, 49, 17, 19], and poisoning attacks, which introduce modifications during training [47, 29]. Although most studies target node classification, techniques such as Projected

Gradient Descent (PGD) also extend to graph classification. Nevertheless, only a few methods explicitly address the graph-level setting. For example, [39] leverage Bayesian optimization to mount attacks on graph classifiers, while [7] evaluate reinforcement learning, gradient-based optimization, and genetic algorithms that iteratively refine graphs via a fitness-driven evolutionary process.

On the defense side, numerous strategies have been proposed to strengthen GNN robustness such as input preprocessing [13]. For instance, [42] remove potentially adversarial edges using Jaccard similarity. These techniques aim to suppress perturbations before message passing. More recent approaches intervene in the message-passing process itself. GNNGuard [46] employs attention-based edge pruning to discard vulnerable connections, and SoftMedian [17] introduces a robust aggregation scheme that down-weights neighbors based on their distance to the dimension-wise median, mitigating the impact of outliers or malicious inputs. Other architectural defenses include GCORN [1], which enforces orthogonal weights to enhance stability. Beyond architectural changes, adversarial training remains highly effective. For example, [19] train GNNs on perturbed graphs, increasing robustness to future attacks. Complementing empirical defenses, certified methods are gaining traction. Techniques such as randomized smoothing [4], building on earlier work [50, 3], offer probabilistic guarantees that predictions remain stable under bounded perturbations.

In parallel, research on pooling strategies has also made great strides [45, 31, 27, 2], where different pooling methods have been presented to enhance the model’s predictive ability. In this work, we aim to connect these two research directions, by examining how flat pooling operations affect the expected robustness. Rather than competing with existing defenses, we offer a complementary viewpoint by analyzing the final pooling stage – a crucial but often overlooked component in robustness research for graph-level prediction. To the best of our knowledge, this is the first systematic study of pooling in the context of adversarial defense for graph classification. While recent work by [40] offers a broad empirical benchmark of hierarchical pooling methods, assessing their robustness and generalization, their study centers on graph coarsening and does not address flat pooling or its theoretical implications in adversarial settings.

3 Preliminaries

In this section, we introduce the fundamental concepts and notation that underpin our work.

Notation and Setup. Let $G = (V, E)$ be a graph, where V and E denote its sets of nodes and edges, respectively. We use $n = |V|$ and $m = |E|$ to represent the number of nodes and edges. The neighborhood of a node $v \in V$ is denoted by $\mathcal{N}(v) = \{u : (v, u) \in E\}$. The degree of node v is thus $\deg(v) = |\mathcal{N}(v)|$. A graph is typically represented by its adjacency matrix $A \in \mathbb{R}^{n \times n}$ and, when available, a node feature matrix $X \in \mathbb{R}^{n \times d}$, where d is the feature dimensionality.

Message Passing GNNs. A GNN consists of a sequence of neighborhood aggregation layers that update node representations by incorporating information from local neighborhoods. Let $h_v^{(0)}$ denote the initial feature vector of node v . The hidden representation of node v at layer $\ell \leq L$ is updated as:

$$a_v^{(\ell)} = \text{AGGREGATE}^{(\ell)}\left(\{h_u^{(\ell-1)} : u \in \mathcal{N}(v)\}\right); \quad h_v^{(\ell)} = \text{COMBINE}^{(\ell)}\left(h_v^{(\ell-1)}, a_v^{(\ell)}\right),$$

where $\text{AGGREGATE}^{(\ell)}$ is a permutation-invariant function that summarizes the features of v ’s neighbors, and $\text{COMBINE}^{(\ell)}$ integrates this summary with v ’s previous representation.

Pooling Operation. After L layers of message passing, a permutation-invariant pooling function is applied to the final node representations to obtain a graph-level embedding:

$$h_G = \text{POOL}\left(\{h_v^{(L)} : v \in V\}\right).$$

This step, often referred to as readout or flat pooling, is essential for producing a fixed-size representation of the entire graph. In this study, we aim to understand how this operation affects the expected robustness of GNNs. For tractability, we focus on flat pooling methods that aggregate node representations in a single step, specifically Sum, Average, and Max pooling.

4 On the Robustness of Pooling Operations

We start by discussing the concept of the expected robustness in graph classification tasks, followed by a theoretical analysis of widely used flat pooling operations.

4.1 Expected Graph Robustness

We consider a set of input graphs represented by their adjacency matrices and node features with associated labels, denoted as $(A_1, X_1, y_1), \dots, (A_N, X_N, y_N) \in (\mathcal{A}, \mathcal{X}, \mathcal{Y})^N$, sampled from an underlying distribution $\mathcal{D}$ over $(\mathcal{A}, \mathcal{X}, \mathcal{Y})$. The goal of graph classification is to learn a function $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ that predicts a target label by minimizing the expected classification risk under $\mathcal{D}$.

In this work, we focus on the black-box evasion setting, where the attacker has no access to the trained model f or its training data and cannot modify them. For the purpose of analysis, we extend the framework introduced in [10] to the graph classification context. Let $(A, X) \in (\mathcal{A}, \mathcal{X})$ be an input graph with label $y \in \mathcal{Y}$. We assess the model’s robustness by examining its behavior in the neighborhood of the input graph, measuring the expected risk of prediction changes under bounded perturbations. Given a perturbation budget ϵ , the expected adversarial risk of f with respect to $\mathcal{D}$ is defined as:

$$\mathcal{R}_\epsilon[f] = \mathbb{E}_{\substack{(A, X) \sim \mathcal{D} \\ (\tilde{A}, \tilde{X}) \in \mathcal{N}_\epsilon(A, X)}} [d_{\mathcal{Y}}(f(\tilde{A}, \tilde{X}), f(A, X))], \quad (1)$$

where $d_{\mathcal{Y}} = \|\cdot\|_2$ denotes distance in the output space $\mathcal{Y}$. The set $\mathcal{N}_\epsilon(A, X) = \{(\tilde{A}, \tilde{X}) : d_{\mathcal{A}, \mathcal{X}}((A, X), (\tilde{A}, \tilde{X})) < \epsilon\}$ defines the neighborhood of admissible perturbations under the budget ϵ , measured by the input space distance $d_{\mathcal{A}, \mathcal{X}}$. In our work, we consider a distance that reflect both the topology and the node features while taking into account the set of permutation matrices Π , which can be written as: $d_{\mathcal{A}, \mathcal{X}}((A, X), (\tilde{A}, \tilde{X})) = \min_{P \in \Pi} \{\|A - P\tilde{A}P^T\|_2 + \|X - P\tilde{X}\|_2\}$.

From a defense perspective, the goal is to minimize the expected risk $\mathcal{R}_\epsilon[f]$, thereby ensuring the model’s predictions remain stable within the allowed perturbation budget. Bounding this risk offers valuable insights into a model’s expected robustness and its sensitivity to adversarial inputs. Based on this, we formalize the notion of GNN robustness in Definition 4.1.

Definition 4.1. (Expected Robustness) A graph-based function $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ is said to be (ϵ, γ) -robust if its expected risk satisfies $\mathcal{R}_\epsilon[f] \leq \gamma$.

We note that this definition is different from the usually used adversarial robustness in the literature, and rather adopts an *average case* perspective by evaluating robustness over the entire neighborhood $\mathcal{N}_\epsilon(A, X)$, in contrast to *worst case* analysis, which considers only the most damaging perturbation. We have chosen to operate within this direction, since we consider that worst case attacks represent a subset of the neighborhood and therefore it follows that upper-bounding this quantity for the average would also give me some insights on the robustness under the worst case criterion. This distinction between average and worst case robustness has been discussed in the literature [36].

4.2 On the Expected Robustness of Flat Pooling Operations

Having established a formal definition of the expected robustness in the context of graph classification, we now apply this framework to assess the robustness properties of commonly used pooling operations. Specifically, we examine two representative models from the message-passing family introduced in Section 3: Graph Convolutional Networks (GCNs) [24] and Graph Isomorphism Networks (GINs) [44], both of which are widely adopted in graph-level tasks.

Following message passing, node embeddings are aggregated through a pooling operation to produce a global graph representation suitable for downstream classification. However, this step can introduce information loss, prompting the development of various pooling strategies. In this work, we focus on three widely used flat pooling techniques, which generate graph-level representations through a single aggregation step over node embeddings: Sum, Average, and Max pooling [9]. While our theoretical analysis centers on these operations, the proposed robustness framework is general and can be extended to other mechanisms. The primary objective is to quantify the influence of each pooling strategy on the expected robustness of the model, as defined in Definition 4.1. Throughout the remainder of this paper, $\|\cdot\|$ denotes the operator norm.

Theorem 4.2. Let $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ denote a graph-based function composed of L GCN layers, where the weight matrix of the ℓ -th layer is denoted by $W^{(\ell)}$. Under a feature-based adversarial attack with perturbation budget ϵ , the robustness of f in the sense of Definition 4.1 satisfies the following:

- If f is based on Max pooling, then it is (ϵ, γ) -robust with:

$$\gamma = \sqrt{\min\{n, d_L\}} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \max_{u \in V} \hat{w}_u \epsilon.$$

- If f is based on Sum pooling, then it is (ϵ, γ) -robust with:

$$\gamma = \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u \in V} \hat{w}_u \epsilon.$$

- If f is based on Average pooling, then it is (ϵ, γ) -robust with:

$$\gamma = \frac{1}{n} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u \in V} \hat{w}_u \epsilon.$$

Here, $\hat{w}_u$ denotes the sum of normalized walks of length $L-1$ originating from node u , and d_L is the output dimensionality of the final layer.

Theorem 4.2 establishes upper bounds on the expected risk for different flat pooling operations. A common factor across all bounds is the product of the weight matrix norms, reflecting the influence of message-passing layers. The second term in each bound differentiates the pooling methods. For Sum pooling, the bound scales with the total normalized walk count, which increases with graph density. Consequently, adversarial noise propagates through multiple paths, weakening robustness and highlighting Sum pooling’s vulnerability in dense graphs. In contrast, Average pooling normalizes this term by the number of nodes, diminishing the impact of increased walk counts. As graphs grow larger or sparser, this normalization tightens the bound, offering improved robustness over Sum pooling. We observe the bound on Max pooling to be governed by a single node in the graph for which the normalized sum of walks of length $L-1$ is maximal. Consequently, the max-pooling operator more localized than the Sum and Mean Operators for which our bounds depend on global characteristics.

The robustness of each pooling method further depends on the attack strategy. Under targeted attacks, where high-degree or influential nodes are perturbed, Max pooling can be particularly vulnerable, especially in gradient-based settings that expose critical nodes. However, when less influential nodes are targeted, it may retain robustness. In non-targeted attacks that perturb nodes more uniformly, Max pooling may outperform Sum and Average pooling, as it aggregates over only one node. Overall, the analysis reveals that no single pooling method is optimal across all scenarios. Robustness varies with graph structure and attack type, underscoring the importance of selecting pooling strategies based on the anticipated threat model. Although our analysis focuses on GCNs, the same theoretical framework naturally extends to Graph Isomorphism Networks (GINs).

Theorem 4.3. Let $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ be composed of L GIN-layers (with its internal parameter $\zeta = 0$) and let $W^{(\ell)}$ denote the weight matrix of the ℓ -th MLP layer. We consider the input node feature space to be bounded, i.e., $\|X\|_2 < B$ for some $B \in \mathbb{R}$. Under a feature-based adversarial attack with perturbation budget ϵ , the model f is (ϵ, γ) -robust with respect to Definition 4.1, with:

- If f is based on Max pooling: $\gamma = \prod_{\ell=1}^L \|W^{(\ell)}\| (BL (\max_{u \in V} \deg(u)) + \epsilon)$.
- If f is based on Sum pooling: $\gamma = \prod_{\ell=1}^L \|W^{(\ell)}\| (2BL|E| + n\epsilon)$.
- If f is based on Average pooling: $\gamma = \sqrt{n, d_L} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \left(\frac{2BL|E|}{n} + \epsilon \right)$.

Here, $|E|$ is the number of edges, n number of nodes and d_L the dimensionality of the final layer L .

As in the GCN case, the derived upper bounds for GIN highlight the influence of pooling operations on expected robustness. For Sum pooling, the expected adversarial risk grows with both the number of nodes and edges, indicating increased vulnerability in larger or denser graphs. In comparison, Average pooling normalizes by the number of nodes, which reduces the relative impact of individual perturbations and enhances robustness in large or sparse settings. Max pooling produces a bound governed by the maximum node degree, suggesting heightened sensitivity in graphs with hub-like structures, particularly when high-degree nodes are targeted. These findings are consistent with GCN-based analysis and further emphasize that pooling is a key factor in shaping the robustness of GNNs.

5 Robust Pooling Through Singular Decomposition

Section 4.2 demonstrated that pooling significantly influences the robustness of graph classification models. This aligns with the intuition behind message passing, where node-level perturbations propagate through local neighborhoods and ultimately affect the global representation, depending on the pooling strategy. Influential nodes with high centrality or strong connectivity tend to amplify adversarial effects, spreading their influence and creating compound vulnerabilities. Similar trends have been observed in node classification [28], where gradient-based attacks often target structurally important nodes. Mitigating the influence of such nodes can therefore improve robustness.

To address adversarial vulnerabilities, prior work has proposed various preprocessing techniques [42, 13, 46, 17]. However, these defenses typically operate before or during the message-passing stage and are primarily tailored for node classification tasks, where graphs are large and removing a subset of nodes has limited impact on performance. In contrast, graph classification often involves smaller graphs for which the removal of nodes can be more critical. Removing or down-weighting certain nodes can disrupt information propagation and degrade clean accuracy, which is especially problematic given that it is not known in advance whether a graph has been attacked. As a result, achieving a favorable trade-off between robustness to adversarial inputs and strong performance on clean data remains a key challenge in graph-level learning.

Given these considerations, to avoid the consequences of node removal, we propose an alternative pooling strategy that effectively balances robustness and clean performance. We introduce *Robust Singular Pooling (RS-Pool)*, a novel method that retains informative signals in clean graphs while mitigating the influence of adversarial perturbations. RS-Pool generates the graph-level representation by extracting the dominant right singular vector of the node embeddings. This design is motivated by classical results in matrix perturbation theory [8, 41], which demonstrate that leading singular vectors are stable under bounded perturbations and encode the most reliable directions of variation.

In the context of graph classification, adversarial attacks distort node embeddings produced by message-passing layers. These distortions are typically localized in nature, due to the limited depth of GNNs (empirically often $L \in \{2, 3, 4, 5\}$ is observed to be highest-performing [14]) and constrained perturbation budgets that restrict modifications to a small subset of nodes [28]. As a result, the dominant singular vector often remains stable, while less-dominant components are more susceptible to noise. By projecting onto this robust direction, RS-Pool generates a stable and expressive graph representation that maintains performance under adversarial conditions. Formally, let $H \in \mathbb{R}^{n \times d}$ denote the matrix of node representations obtained after the message-passing stage (e.g., from a GCN [24]), where n is the number of nodes and d the features (embedding) dimension. We consider the singular value decomposition (SVD): $H = U\Sigma V^\top$, where $U \in \mathbb{R}^{n \times n}$, $\Sigma \in \mathbb{R}^{n \times d}$, and $V \in \mathbb{R}^{d \times d}$. The singular values $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_{\min(n,d)} \geq 0$ quantify the contribution of each singular direction. Let $v_1 \in \mathbb{R}^d$ denote the dominant right singular vector, RS-Pool defines the graph-level representation as a scaled version of v_1 :

$$\begin{aligned} \text{RS-Pool: } \quad \mathbb{R}^{n \times d} &\rightarrow \mathbb{R}^d \\ H &\mapsto \tau v_1(H), \end{aligned}$$

where $\tau \in \mathbb{R}_{>0}$ is a scaling factor controlling the magnitude of the output embedding. The dominant singular vector captures the principal direction of variation in the embedding space, effectively summarizing the global structure encoded in H . Scaling ensures the pooled vector retains sufficient energy to serve as a meaningful input to downstream classifiers.

Crucially, RS-Pool also offers inherent robustness. Classical results from matrix perturbation theory [8, 41] show that if the leading singular value $\sigma_1(H)$ is well-separated from the next $\sigma_2(H)$ (we use σ_s and $\sigma_s(H)$ interchangeably when unambiguous), then the top singular vector $v_1(H)$ is stable under small perturbations to H . This separation condition ensures that adversarial distortions, often low-rank and localized, have limited influence on the dominant representation. Within the expected adversarial risk framework introduced in Definition 4.1, we now formalize the robustness guarantee of RS-Pool.

Theorem 5.1. *Let $f: (\mathcal{G}, \mathcal{X}) \rightarrow \mathcal{Y}$ denote a graph function composed of L GCN layers and using our RS-Pool, where the weight matrix of the i -th layer is denoted by $W^{(i)}$. Under a feature-based*

adversarial attack with perturbation budget ϵ , f is (ϵ, γ) -robust in respect to Definition 4.1 with:

$$\gamma = \frac{\tau\sqrt{2}\epsilon}{\sigma_1 - \sigma_2} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u=1}^n (\hat{w}_u)^2, \quad (2)$$

with $\hat{w}_u$ denoting the sum of normalized walks of length $(L - 1)$ starting from node u , and $\sigma_1 \neq \sigma_2$ being the two dominant singular values.

Theorem 5.1 provides an upper bound on the expected adversarial risk of the proposed RS-Pool method. As with previous bounds, the expression includes shared terms such as the product of message-passing weight norms and the perturbation budget. However, a distinguishing feature of this bound is its dependence on the spectral gap of the clean representation matrix H and the scaling factor τ . Notably, the bound relies solely on the spectral gap of the unperturbed node embeddings, not the perturbed input, reinforcing our earlier claim that RS-Pool mitigates the influence of adversarial noise. Furthermore, the scaling parameter τ acts as a tunable control for the representation’s sensitivity, allowing the bound to be tightened and thereby improving robustness. This tunability is especially beneficial in settings where the message-passing layers exhibit high sensitivity to input perturbations.

Corollary 5.2. *Let $f: (\mathcal{G}, \mathcal{X}) \rightarrow \mathcal{Y}$ be the graph-based function considered in Theorem 5.1. Let γ be the upper-bound derived in Theorem 5.1, we have that f is (ϵ, γ') -robust with:*

$$\gamma' = \min\{\gamma, 2\tau\}. \quad (3)$$

Corollary 5.2 highlights that even when the message-passing component is highly vulnerable to adversarial perturbations, as indicated by large weight norms, RS-Pool can still ensure a bounded expected risk by constraining the representation through the scaling parameter τ . This scenario has motivated prior robustness techniques such as orthogonal regularization [6]. The result demonstrates that RS-Pool serves as a stabilizing mechanism, contributing to overall robustness even when earlier layers are sensitive to input perturbations. Moreover, RS-Pool retains the essential property of permutation invariance, ensuring its compatibility with standard GNN architectures.

Lemma 5.3. *RS-Pool is permutation invariant; that is, for any pair of isomorphic graphs G and G_Π with its corresponding permutation matrix Π , we have $f(A, X) = f(A_\Pi, X_\Pi)$.*

Remark. Theorem 5.1 does not address cases where the largest singular value has multiplicity greater than one. Such cases are rare in practice, as GNNs’ inherent smoothing tends to reduce the feature matrix rank and create a non-trivial spectral gap. Nonetheless, Wedin’s Theorem, which underlies our proof, extends to the case where the dominant singular value has multiplicity r , yielding an additional factor $\sqrt{r}$ (instead of $\sqrt{2}$) in the upper bound γ , with the relevant spectral gap $\sigma_r - \sigma_{r+1}$.

Implementation and Algorithm. RS-Pool constructs the graph-level representation by extracting the dominant right-singular vector of the node embedding matrix H . As only the top singular vector is needed, computing the costly full SVD is unnecessary. Instead, we employ the truncated power iteration method to efficiently estimate the dominant vector. This iterative approach is scalable with respect to both the number of nodes and the embedding dimension. A key advantage of RS-Pool is that it remains fully differentiable, allowing seamless integration into existing GNN pipelines and enabling end-to-end training via standard backpropagation. Algorithm 1 outlines the forward pass of a GNN model utilizing RS-Pool.

Algorithm 1 RS-Pool Forward Pass

Require: $H \in \mathbb{R}^{n \times d}$, $\tau \in \mathbb{R}_{>0}$, $K \in \mathbb{N}_{>0}$
1: $S \leftarrow H^\top H$
2: $v \leftarrow$ random unit vector in $\mathbb{R}^d$
3: **for** $t = 1$ **to** K **do**
4: $v \leftarrow Sv$
5: $v \leftarrow v / \|v\|_2$
6: **end for**
7: **return** τHv

Computational Complexity. The main computational overhead of RS-Pool stems from approximating the dominant singular vector using power iteration. Each iteration has a complexity of $\mathcal{O}(n \times d)$, where n is the number of nodes and d the embedding dimension. In practice, we find that only a small number of iterations (e.g., 2–5) is sufficient to obtain a reliable estimate. This efficiency is supported by the convergence properties of power iteration, which is known to converge geometrically at a rate determined by the ratio $\sigma_2(H)/\sigma_1(H)$, where $\sigma_1(H)$ and $\sigma_2(H)$ are the top two singular values of H . A well-separated spectral gap ensures rapid convergence, which is commonly observed in GNN embeddings due to dominant low-frequency components. As a result, the overall complexity remains $\mathcal{O}(K \times n \times d)$ for a small K , making RS-Pool suitable for real-world graph applications. A comparison of its runtime with other pooling methods is provided in Appendix E.2.

We further note that the fast empirical convergence of power iteration aligns with the theoretical robustness bound in Theorem 5.1, where the spectral gap $\sigma_1(H) - \sigma_2(H)$ plays a central role. A larger gap not only tightens the bound (lower γ), indicating improved robustness, but also ensures faster convergence of the power iteration method used to compute RS-Pool. This dual benefit highlights the practical effectiveness of RS-Pool, as favorable spectral properties commonly observed in GNN embeddings lead to both stable representations and efficient computation. We provide empirical validation of this observation in Appendix E.1.

6 Experimental Results

This section aims to validate our theoretical insights on real-world datasets where we aim to investigate the robustness of RS-Pool in comparison to other considered pooling methods. In line with our theoretical study, we focus on the GCN model in this section while we provide results on the GIN in Appendix E.4. We additionally provide results when considering the adversarial defenses as a backbone GNN in Appendix E.5, showcasing therefore the universality of our insights. For all our experiments, we use a 2 layers GCN (resp. GIN) classifier with identical hyper-parameters across all experiments to guarantee a fair comparison. Each experiment is run 10 times to reduce the effect of randomness. Additional implementation details are provided in Appendix F.

Considered Benchmarks. We evaluate RS-Pool against classical flat pooling methods, including Average, Max, and Sum pooling, as well as several advanced approaches. We include Self-Attention Graph pooling (SAG) [27], which uses node features and graph structure to compute attention scores and retain informative nodes, and TopK pooling (TopK-P) [16], which learns to rank nodes via a projection vector and selects a fixed fraction by score. We also consider Path-Integral-based Pooling (PAN-P) [31], a TopK extension that integrates multi-hop structure into scoring. Lastly, we compare with Sort Pooling (Sort-P) [45], which orders nodes by feature magnitude and concatenates the top- k sorted nodes.

Attacks. We evaluate robustness under three adversarial attack strategies: (i) *Random Attack*, which randomly adds or removes edges in the input graph. The search is performed over K random perturbations, and the one yielding the worst performance is selected; (ii) *Genetic Attack* [7], which uses evolutionary algorithms to generate adversarial graphs via selection, crossover, and mutation; and (iii) *Gradient-Based Attack (PGD)* [7], which greedily modifies the graph structure by targeting edges with the highest gradient magnitude relative to the model’s input. For all attacks, we apply a perturbation budget of $\epsilon = 0.3$, allowing up to 30% of the edges in each graph to be modified. We additionally considered the Bit-Flip Attacks (BFAs) [26, 25] in Appendix E.6.

Datasets. We conduct experiments on standard graph classification datasets from the TUDataset benchmark [32], spanning diverse domains. In bioinformatics graphs (PROTEINS, D&D, ENZYMES), small changes to residue contact links can influence biological property predictions. In molecular graphs (NCI1, ER_MD), altering a bond may change the predicted molecular function. For social networks (IMDB-B, REDDIT-B), edge perturbations, such as fake user interactions, can flip the predicted graph label. In image-based graphs (MSRC_9), local edits to image patches can lead to misclassification. For all datasets, we use the public train/validation/test splits when available [14]; otherwise, we adopt the same evaluation protocol and report the specific folds used.

6.1 Experimental Results

Empirical Estimation of Adversarial Risk. We begin by empirically analyzing the adversarial risk defined in Section 4, characterized by the upper bound γ . To this end, we compute the distance between the pooled representations of clean and adversarially perturbed graphs for each pooling method. Figure 1 presents the results of this analysis across the considered pooling strategies. As expected, RS-Pool consistently maintains a smaller distance between clean and perturbed representations. In addition, we see the exact behavior underlined in our theoretical results provided in Theorem 4.2, where the Sum is the one with the higher upper-bound, followed by the Max and the Average. This behavior directly correlates with lower attack success rates of our RS-Pool, as shown in subplots (b) and (d). Complete results for all pooling methods are provided in Appendix E.3.

On the Effect of the Hyperparameter τ . As discussed in Section 5, the parameter τ plays a critical role in RS-Pool. While it directly scales the extracted dominant singular vector, its more significant

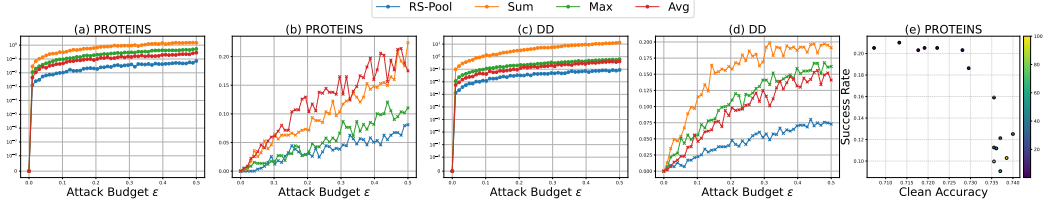


Figure 1: Left: Empirical estimation of adversarial risk γ and corresponding attack success rate on PROTEINS ((a), (b)) and D&D ((c), (d)). Right (e): Study of the effect of the parameter α which inversely controls τ on the PROTEINS dataset.

Table 1: Clean and Attacked classification accuracy ($\pm$ standard deviation) of the RS-Pool and other pooling strategies on different graph classification dataset when subject to adversarial attacks.

Dataset	Attack	Sum	Average	Max	SAG	TopK-P	PAN-P	Sort-P	RS-Pool
PROTEINS	Clean	74.2 $\pm$ 3.1	70.8 $\pm$ 2.2	73.2 $\pm$ 2.3	72.3 $\pm$ 3.3	73.2 $\pm$ 2.9	74.7 $\pm$ 1.7	72.3 $\pm$ 2.3	73.5 $\pm$ 2.9
	PGD	45.8 $\pm$ 2.9	34.2 $\pm$ 2.1	28.2 $\pm$ 3.5	45.5 $\pm$ 2.1	48.5 $\pm$ 1.2	33.3 $\pm$ 1.5	31.8 $\pm$ 2.9	51.9 $\pm$ 3.6
	Random	68.3 $\pm$ 3.1	65.7 $\pm$ 2.3	42.9 $\pm$ 4.1	61.9 $\pm$ 5.1	68.2 $\pm$ 1.8	68.5 $\pm$ 2.1	64.3 $\pm$ 3.2	70.4 $\pm$ 3.1
	Genetic	66.1 $\pm$ 2.3	65.3 $\pm$ 3.2	39.8 $\pm$ 5.0	62.9 $\pm$ 3.8	66.5 $\pm$ 2.2	66.4 $\pm$ 2.7	63.7 $\pm$ 2.1	68.2 $\pm$ 1.8
D&D	Clean	75.1 $\pm$ 0.6	70.1 $\pm$ 0.5	74.1 $\pm$ 0.6	71.3 $\pm$ 0.4	74.4 $\pm$ 1.1	75.7 $\pm$ 2.8	74.9 $\pm$ 1.7	74.6 $\pm$ 0.7
	PGD	6.7 $\pm$ 2.0	12.6 $\pm$ 3.6	8.7 $\pm$ 5.2	9.0 $\pm$ 3.8	5.9 $\pm$ 2.9	9.6 $\pm$ 2.8	16.6 $\pm$ 1.9	30.4 $\pm$ 3.2
	Random	61.9 $\pm$ 3.4	64.4 $\pm$ 2.9	19.7 $\pm$ 2.4	62.8 $\pm$ 2.4	58.6 $\pm$ 1.1	61.9 $\pm$ 2.4	50.3 $\pm$ 2.9	66.7 $\pm$ 2.1
	Genetic	60.5 $\pm$ 4.2	63.1 $\pm$ 2.6	21.4 $\pm$ 1.5	62.5 $\pm$ 2.8	57.7 $\pm$ 0.6	56.1 $\pm$ 3.6	57.7 $\pm$ 4.3	67.2 $\pm$ 2.9
NCH	Clean	70.6 $\pm$ 0.8	67.9 $\pm$ 1.6	68.2 $\pm$ 1.2	70.4 $\pm$ 0.9	70.9 $\pm$ 0.9	70.4 $\pm$ 1.7	69.5 $\pm$ 0.4	70.1 $\pm$ 1.2
	PGD	23.4 $\pm$ 1.1	14.1 $\pm$ 2.1	19.7 $\pm$ 1.7	22.3 $\pm$ 1.3	23.2 $\pm$ 2.7	23.9 $\pm$ 2.3	23.5 $\pm$ 1.1	26.3 $\pm$ 1.8
	Random	26.1 $\pm$ 0.7	20.6 $\pm$ 2.6	9.6 $\pm$ 5.1	19.8 $\pm$ 1.5	25.8 $\pm$ 1.8	24.7 $\pm$ 2.1	24.3 $\pm$ 1.2	27.2 $\pm$ 0.1
	Genetic	25.9 $\pm$ 0.8	19.7 $\pm$ 2.7	9.3 $\pm$ 2.2	18.7 $\pm$ 1.2	23.9 $\pm$ 2.4	19.1 $\pm$ 2.6	25.8 $\pm$ 1.4	26.9 $\pm$ 0.7
ENZYMES	Clean	33.4 $\pm$ 4.9	27.7 $\pm$ 4.9	27.7 $\pm$ 1.6	26.1 $\pm$ 0.9	27.7 $\pm$ 4.3	28.3 $\pm$ 3.6	29.7 $\pm$ 2.1	32.8 $\pm$ 4.6
	PGD	2.7 $\pm$ 2.1	3.3 $\pm$ 1.3	9.8 $\pm$ 2.3	8.9 $\pm$ 2.1	6.7 $\pm$ 1.4	1.7 $\pm$ 1.3	5.6 $\pm$ 2.1	11.9 $\pm$ 2.7
	Random	7.4 $\pm$ 3.4	8.3 $\pm$ 2.3	2.2 $\pm$ 1.6	3.8 $\pm$ 1.6	10.6 $\pm$ 1.3	5.6 $\pm$ 2.1	9.4 $\pm$ 1.6	11.7 $\pm$ 2.1
	Genetic	8.9 $\pm$ 4.3	7.2 $\pm$ 4.1	6.1 $\pm$ 0.8	4.4 $\pm$ 2.1	9.4 $\pm$ 1.6	5.0 $\pm$ 1.3	9.4 $\pm$ 1.6	9.9 $\pm$ 2.1
IMDB-B	Clean	52.9 $\pm$ 4.2	62.7 $\pm$ 3.9	56.7 $\pm$ 4.1	63.7 $\pm$ 4.1	52.0 $\pm$ 5.3	53.3 $\pm$ 2.4	56.7 $\pm$ 4.2	62.0 $\pm$ 3.2
	PGD	52.0 $\pm$ 2.8	61.3 $\pm$ 4.4	32.0 $\pm$ 5.8	61.3 $\pm$ 3.2	47.0 $\pm$ 4.5	49.6 $\pm$ 2.9	56.7 $\pm$ 3.2	61.7 $\pm$ 2.5
	Random	52.3 $\pm$ 3.2	62.3 $\pm$ 3.2	45.9 $\pm$ 3.3	62.9 $\pm$ 4.3	49.7 $\pm$ 4.0	51.3 $\pm$ 1.2	55.7 $\pm$ 3.3	61.3 $\pm$ 2.9
	Genetic	52.0 $\pm$ 2.8	61.0 $\pm$ 2.3	52.0 $\pm$ 4.1	61.7 $\pm$ 3.8	50.9 $\pm$ 4.3	52.3 $\pm$ 1.7	57.8 $\pm$ 3.3	62.7 $\pm$ 3.2
REDDIT-B	Clean	77.5 $\pm$ 0.8	70.0 $\pm$ 1.2	72.0 $\pm$ 0.9	70.5 $\pm$ 1.6	78.5 $\pm$ 1.4	77.5 $\pm$ 1.5	80.0 $\pm$ 0.9	75.5 $\pm$ 0.6
	PGD	69.0 $\pm$ 1.4	57.5 $\pm$ 2.5	62.0 $\pm$ 2.5	55.5 $\pm$ 1.9	55.5 $\pm$ 1.9	57.0 $\pm$ 1.9	56.5 $\pm$ 2.5	74.0 $\pm$ 0.9
	Random	76.0 $\pm$ 0.8	69.5 $\pm$ 0.9	67.9 $\pm$ 1.2	69.0 $\pm$ 1.6	76.5 $\pm$ 1.5	75.6 $\pm$ 1.4	78.0 $\pm$ 0.6	75.0 $\pm$ 0.3
	Genetic	74.0 $\pm$ 0.9	69.0 $\pm$ 1.5	66.9 $\pm$ 1.4	68.5 $\pm$ 1.9	55.5 $\pm$ 1.5	57.0 $\pm$ 2.5	56.5 $\pm$ 1.9	74.5 $\pm$ 0.7
ER_MD	Clean	64.0 $\pm$ 3.9	61.1 $\pm$ 6.7	61.4 $\pm$ 6.1	61.1 $\pm$ 5.8	66.3 $\pm$ 4.2	69.7 $\pm$ 1.8	67.1 $\pm$ 5.4	65.9 $\pm$ 5.9
	PGD	40.4 $\pm$ 3.2	44.5 $\pm$ 3.8	45.8 $\pm$ 3.3	45.9 $\pm$ 2.8	37.4 $\pm$ 5.3	35.2 $\pm$ 4.5	34.1 $\pm$ 6.2	47.9 $\pm$ 3.4
	Random	55.1 $\pm$ 2.3	51.6 $\pm$ 2.8	52.4 $\pm$ 3.8	49.8 $\pm$ 2.3	50.1 $\pm$ 3.8	55.1 $\pm$ 2.3	52.4 $\pm$ 3.8	58.1 $\pm$ 5.9
	Genetic	60.2 $\pm$ 1.1	58.9 $\pm$ 1.8	60.2 $\pm$ 0.8	59.9 $\pm$ 2.3	51.4 $\pm$ 1.8	52.4 $\pm$ 4.2	55.1 $\pm$ 2.3	61.7 $\pm$ 1.3
MSRC_9	Clean	88.6 $\pm$ 3.2	88.6 $\pm$ 1.8	88.6 $\pm$ 1.8	90.1 $\pm$ 1.1	87.1 $\pm$ 2.8	89.4 $\pm$ 2.1	89.4 $\pm$ 2.8	89.7 $\pm$ 1.1
	PGD	82.6 $\pm$ 2.8	84.1 $\pm$ 2.1	78.0 $\pm$ 2.9	82.6 $\pm$ 2.3	54.6 $\pm$ 1.9	84.1 $\pm$ 2.3	75.8 $\pm$ 3.4	85.6 $\pm$ 3.1
	Random	85.6 $\pm$ 2.6	87.8 $\pm$ 1.3	86.3 $\pm$ 1.3	87.1 $\pm$ 1.6	84.1 $\pm$ 3.2	86.8 $\pm$ 1.3	86.4 $\pm$ 2.1	89.1 $\pm$ 0.4
	Genetic	85.6 $\pm$ 2.1	86.3 $\pm$ 2.8	87.1 $\pm$ 1.9	87.1 $\pm$ 2.1	84.1 $\pm$ 3.2	86.3 $\pm$ 1.8	86.4 $\pm$ 3.2	89.4 $\pm$ 1.1

effect lies in modulating the robustness bound in Theorem 5.1 and influencing the expected risk as detailed in Corollary 5.2. To validate this theoretically grounded effect, we examine how varying τ impacts both clean accuracy and robustness. In our experiments, to better isolate the role of the spectral gap, we parameterize τ as a scaled version of the leading singular value, which is already computed during the power iteration. Specifically, we set $\tau = \sigma_1(X)/\alpha$, where $\alpha > 0$ is a tunable constant. We then study how different values of α affect performance, measuring clean accuracy and attack success rate. As shown in Figure 1 (e), increasing α (and thus decreasing τ) leads to a lower attack success rate, indicating improved robustness. However, this comes at the cost of clean accuracy, which peaks at an intermediate α before declining. These results highlight the importance of selecting an appropriate value of α to achieve a desirable trade-off between robustness and predictive performance.

Results Analysis. We evaluate the performance of RS-Pool and baseline pooling methods under adversarial attacks. Table 1 presents classification accuracy on clean and perturbed data across datasets and attack types. RS-Pool consistently performs well on clean inputs and significantly outperforms baselines under adversarial conditions. The only exceptions occur in two cases involving the Random attack, where the limited perturbation budget results in low attack success, allowing some baselines to retain higher accuracy due to strong clean performance. Under stronger attacks

Table 2: Clean and Attacked classification accuracy ($\pm$ standard deviation) of the RS-Pool and other pooling strategies on different graph classification dataset when subject to feautre-based adversarial attacks.

Dataset	Attack	Sum	Average	Max	SAG	TopK-P	PAN-P	Sort-P	RS-Pool
PROTEINS	Clean	74.2 $\pm$ 3.1	70.8 $\pm$ 2.2	73.2 $\pm$ 2.3	72.3 $\pm$ 3.3	73.2 $\pm$ 2.9	74.7 $\pm$ 1.7	72.3 $\pm$ 2.3	73.5 $\pm$ 2.9
	PGD	68.4 $\pm$ 4.0	64.6 $\pm$ 2.2	66.4 $\pm$ 4.1	64.3 $\pm$ 3.6	64.9 $\pm$ 6.9	68.8 $\pm$ 2.9	66.7 $\pm$ 2.9	69.9 $\pm$ 4.1
	Random	69.5 $\pm$ 5.2	67.6 $\pm$ 2.9	68.4 $\pm$ 3.2	69.0 $\pm$ 4.0	68.8 $\pm$ 6.5	69.2 $\pm$ 3.2	70.2 $\pm$ 2.6	70.9 $\pm$ 2.8
D&D	Clean	75.1 $\pm$ 0.6	70.1 $\pm$ 0.5	74.1 $\pm$ 0.6	71.3 $\pm$ 0.4	74.4 $\pm$ 1.1	75.7 $\pm$ 2.8	74.9 $\pm$ 1.7	74.6 $\pm$ 0.7
	PGD	36.9 $\pm$ 3.6	36.9 $\pm$ 6.8	46.4 $\pm$ 6.4	49.9 $\pm$ 1.8	28.7 $\pm$ 3.7	37.5 $\pm$ 4.1	47.9 $\pm$ 4.8	51.8 $\pm$ 1.1
	Random	49.0 $\pm$ 5.2	44.0 $\pm$ 6.5	45.7 $\pm$ 4.6	54.1 $\pm$ 3.0	40.9 $\pm$ 5.1	51.0 $\pm$ 4.6	49.7 $\pm$ 4.9	54.9 $\pm$ 2.3
NCI1	Clean	70.6 $\pm$ 0.8	67.9 $\pm$ 1.6	68.2 $\pm$ 1.2	70.4 $\pm$ 0.9	70.9 $\pm$ 0.9	70.4 $\pm$ 1.7	69.5 $\pm$ 0.4	70.1 $\pm$ 1.2
	PGD	38.6 $\pm$ 2.8	36.4 $\pm$ 3.5	37.9 $\pm$ 5.3	36.6 $\pm$ 2.2	35.1 $\pm$ 2.2	38.4 $\pm$ 1.8	36.2 $\pm$ 3.6	39.4 $\pm$ 1.9
	Random	41.1 $\pm$ 2.0	38.1 $\pm$ 2.9	40.9 $\pm$ 5.5	40.1 $\pm$ 2.9	38.5 $\pm$ 1.0	40.9 $\pm$ 2.4	38.5 $\pm$ 4.8	41.7 $\pm$ 1.3
ER_MD	Clean	64.0 $\pm$ 3.9	61.1 $\pm$ 6.7	61.4 $\pm$ 6.1	61.1 $\pm$ 5.8	66.3 $\pm$ 4.2	69.7 $\pm$ 1.8	67.1 $\pm$ 5.4	65.9 $\pm$ 5.9
	PGD	41.9 $\pm$ 3.2	45.3 $\pm$ 1.4	45.3 $\pm$ 4.6	45.7 $\pm$ 3.5	45.3 $\pm$ 3.8	45.3 $\pm$ 2.1	43.1 $\pm$ 3.5	55.1 $\pm$ 4.9
	Random	54.3 $\pm$ 2.3	50.2 $\pm$ 4.3	52.1 $\pm$ 6.4	53.9 $\pm$ 5.5	52.1 $\pm$ 4.1	51.3 $\pm$ 5.2	53.6 $\pm$ 7.7	56.2 $\pm$ 5.1
MSRC_9	Clean	88.6 $\pm$ 3.2	88.6 $\pm$ 1.8	88.6 $\pm$ 1.8	90.1 $\pm$ 1.1	87.1 $\pm$ 2.8	89.4 $\pm$ 2.1	89.4 $\pm$ 2.8	89.7 $\pm$ 1.1
	PGD	84.8 $\pm$ 2.8	85.6 $\pm$ 4.0	85.6 $\pm$ 2.8	87.1 $\pm$ 1.1	83.3 $\pm$ 1.1	85.6 $\pm$ 3.9	84.1 $\pm$ 1.9	88.3 $\pm$ 1.1
	Random	86.4 $\pm$ 3.7	88.6 $\pm$ 1.9	87.9 $\pm$ 2.8	89.4 $\pm$ 1.1	87.1 $\pm$ 2.1	87.9 $\pm$ 3.9	88.6 $\pm$ 3.2	89.3 $\pm$ 1.1

such as PGD, RS-Pool shows notable gains, achieving up to a 14% uplift over the second-best method on the D&D dataset. This advantage is more pronounced on datasets like D&D, which contain larger and denser graphs where the leading singular vector captures most of the informative signal, enabling more robust representation through RS-Pool. Similar observations hold for GIN and other adversarial defenses, as shown in Appendix E.4 and E.5.

Feature-Based Robustness. To further demonstrate the value of our theoretical analysis, and inline with the considered setup, we evaluate adversarial attacks targeting node features. We adopt the same attack strategies used in the structural setting: *Random*, which adds noise to features at random, and *PGD*, a gradient-based attack. Table 2 reports clean and attacked accuracies (with standard deviations) for our proposed pooling method, RS-Pool, alongside other baseline pooling strategies across multiple datasets. Consistent with the structural attack results, RS-Pool achieves higher attacked accuracy, highlighting its improved robustness against feature perturbations.

7 Conclusion

In this work, we provided a theoretical analysis of how pooling operations affect the adversarial robustness of GNNs in the context of graph classification. Our findings indicate that different pooling strategies exhibit varying levels of sensitivity to attack types – some being more robust to targeted perturbations, while others are more vulnerable under global attacks. To address this, we introduced *RS-Pool*, a novel pooling method that uses a scaled version of the dominant singular vector of the node embedding matrix as the graph-level representation. We theoretically demonstrated that RS-Pool can attenuate the impact of adversarial perturbations at the pooling stage, thereby improving the model’s overall robustness. Empirical evaluations on real-world datasets, using different GNN backbones and adversarial attacks, confirm that RS-Pool achieves stronger robustness while maintaining competitive performance on clean data.

Limitations. Our method relies on the existence of a spectral gap in feature matrices, yet gap size can depend on the network depth and graph topology and varies across settings, which warrants further study. Moreover, power iteration pooling scales linearly with edges but may incur nontrivial overhead on very dense graphs, limiting its use in resource constrained applications.

Acknowledgements

The computation (on GPUs) was partially enabled by resources provided by the National Academic Infrastructure for Supercomputing in Sweden (NAISS) at Alvis partially funded by the Swedish Research Council through grant agreement no. “2024/22-309”. J.L. would additionally like to aknowlegde support by the French National Research Agency (ANR) via the “GraspGNNs” JCJC grant (ANR-24-CE23-3888). We furthermore want to acknowledge all the reviewers for their feedback that improved our manuscript.

References

- [1] Yassine Abbahaddou, Sofiane Ennadir, Johannes F Lutzeyer, Michalis Vazirgiannis, and Henrik Boström. Bounding the expected robustness of graph neural networks subject to node feature attacks. In *International Conference on Learning Representations*, 2024.
- [2] Filippo Maria Bianchi, Daniele Grattarola, and Cesare Alippi. Spectral clustering with graph neural networks for graph pooling. In *International conference on machine learning*, pages 874–883. PMLR, 2020.
- [3] Aleksandar Bojchevski and Stephan Günnemann. Certifiable robustness to graph perturbations, 2019.
- [4] Aleksandar Bojchevski, Johannes Klicpera, and Stephan Günnemann. Efficient robustness certificates for discrete data: Sparsity-aware randomized smoothing for graphs, images and more, 2020.
- [5] Chen Cai, Dingkan Wang, and Yusu Wang. Graph coarsening with neural networks. *arXiv preprint arXiv:2102.01350*, 2021.
- [6] Moustapha Cisse, Piotr Bojanowski, Edouard Grave, Yann Dauphin, and Nicolas Usunier. Parseval networks: Improving robustness to adversarial examples. In *International conference on machine learning*, pages 854–863. PMLR, 2017.
- [7] Hanjun Dai, Hui Li, Tian Tian, Xin Huang, Lin Wang, Jun Zhu, and Le Song. Adversarial Attack on Graph Structured Data. In *Proceedings of the 35th International Conference on Machine Learning*, pages 1115–1124, 2018.
- [8] Chandler Davis and W. M. Kahan. The rotation of eigenvectors by a perturbation. iii. *SIAM Journal on Numerical Analysis*, 7(1):1–46, 1970.
- [9] David K Duvenaud, Dougal Maclaurin, Jorge Iparraguirre, Rafael Bombarell, Timothy Hirzel, Alan Aspuru-Guzik, and Ryan P Adams. Convolutional networks on graphs for learning molecular fingerprints. In C. Cortes, N. Lawrence, D. Lee, M. Sugiyama, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc., 2015.
- [10] Sofiane Ennadir, Yassine Abbahaddou, Johannes F Lutzeyer, Michalis Vazirgiannis, and Henrik Boström. A simple and yet fairly effective defense for graph neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 21063–21071, 2024.
- [11] Sofiane Ennadir, Amr Alkhatib, Giannis Nikolentzos, Michalis Vazirgiannis, and Henrik Boström. Unboundattack: Generating unbounded adversarial attacks to graph neural networks. In *International Conference on Complex Networks and Their Applications*, pages 100–111. Springer, 2023.
- [12] Sofiane Ennadir, Johannes Lutzeyer, Michalis Vazirgiannis, and El Houcine Bergou. If you want to be robust, be wary of initialization. *Advances in Neural Information Processing Systems*, 37:23796–23823, 2024.
- [13] Negin Entezari, Saba A Al-Sayouri, Amirali Darvishzadeh, and Evangelos E Papalexakis. All you need is low (rank) defending against adversarial attacks on graphs. In *Proceedings of the 13th international conference on web search and data mining*, pages 169–177, 2020.
- [14] Federico Errica, Marco Podda, Davide Bacciu, and Alessio Micheli. A Fair Comparison of Graph Neural Networks for Graph Classification. In *8th International Conference on Learning Representations*, 2020.
- [15] Matthias Fey and Jan E. Lenssen. Fast graph representation learning with PyTorch Geometric. In *ICLR Workshop on Representation Learning on Graphs and Manifolds*, 2019.
- [16] Hongyang Gao and Shuiwang Ji. Graph u-nets. In *international conference on machine learning*, pages 2083–2092. PMLR, 2019.

- [17] Simon Geisler, Tobias Schmidt, Hakan Şirin, Daniel Zügner, Aleksandar Bojchevski, and Stephan Günnemann. Robustness of graph neural networks at scale. *Advances in Neural Information Processing Systems*, 34:7637–7649, 2021.
- [18] Justin Gilmer, Samuel S Schoenholz, Patrick F Riley, Oriol Vinyals, and George E Dahl. Neural message passing for quantum chemistry. In *International Conference on Machine Learning*, pages 1263–1272. PMLR, 2017.
- [19] Lukas Gosch, Simon Geisler, Daniel Sturm, Bertrand Charpentier, Daniel Zügner, and Stephan Günnemann. Adversarial training for graph neural networks: Pitfalls, solutions, and new directions. *Advances in Neural Information Processing Systems*, 36, 2024.
- [20] Stephan Günnemann. Graph neural networks: Adversarial robustness. In *Graph Neural Networks: Foundations, Frontiers, and Applications*, pages 149–176. Springer, 2022.
- [21] Wei Jin, Yaxing Li, Han Xu, Yiqi Wang, Shuiwang Ji, Charu Aggarwal, and Jiliang Tang. Adversarial attacks and defenses on graphs. *SIGKDD Explor. Newsl.*, page 19–34, 2021.
- [22] Steven Kearnes, Kevin McCloskey, Marc Berndl, Vijay Pande, and Patrick Riley. Molecular graph convolutions: moving beyond fingerprints. *Journal of Computer-Aided Molecular Design*, 30(8):595–608, 2016.
- [23] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [24] Thomas N. Kipf and Max Welling. Semi-Supervised Classification with Graph Convolutional Networks. In *International Conference on Learning Representations (ICLR)*, 2017.
- [25] Lorenz Kummer, Wilfried N. Gansterer, and Nils Morten Kriege. On the relationship between robustness and expressivity of graph neural networks. In *The 28th International Conference on Artificial Intelligence and Statistics*, 2025.
- [26] Lorenz Kummer, Samir Moustafa, Sebastian Schrittwieser, Wilfried Gansterer, and Nils Kriege. Attacking graph neural networks with bit flips: Weisfeiler and leman go indifferent. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD ’24*, page 1428–1439, New York, NY, USA, 2024. Association for Computing Machinery.
- [27] Junhyun Lee, Inyeop Lee, and Jaewoo Kang. Self-attention graph pooling. In *International conference on machine learning*, pages 3734–3743. pmlr, 2019.
- [28] Kuan Li, Yang Liu, Xiang Ao, and Qing He. Revisiting graph adversarial attack and defense from a data distribution perspective. In *The Eleventh International Conference on Learning Representations*, 2023.
- [29] Vijay Lingam, Mohammad Sadegh Akhondzadeh, and Aleksandar Bojchevski. Rethinking label poisoning for GNNs: Pitfalls and attacks. In *The Twelfth International Conference on Learning Representations*, 2024.
- [30] Chuang Liu, Yibing Zhan, Jia Wu, Chang Li, Bo Du, Wenbin Hu, Tongliang Liu, and Dacheng Tao. Graph pooling for graph neural networks: progress, challenges, and opportunities. In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence, IJCAI ’23*, 2023.
- [31] Zheng Ma, Junyu Xuan, Yu Guang Wang, Ming Li, and Pietro Liò. Path integral based convolution and pooling for graph neural networks. *Advances in Neural Information Processing Systems*, 33:16421–16433, 2020.
- [32] Christopher Morris, Nils M Kriege, Franka Bause, Kristian Kersting, Petra Mutzel, and Marion Neumann. Tudataset: A collection of benchmark datasets for learning with graphs. *arXiv preprint arXiv:2007.08663*, 2020.

- [33] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. Pytorch: An imperative style, high-performance deep learning library. In *Advances in Neural Information Processing Systems* 32, pages 8024–8035. Curran Associates, Inc., 2019.
- [34] Hao Peng, Zehao Yu, Dandan Zhao, Zhiguo Ding, Jieshuai Yang, Bo Zhang, Jianming Han, Xuhong Zhang, Shouling Ji, and Ming Zhong. Evading control flow graph based GNN malware detectors via active opcode insertion method with maliciousness preserving. *Scientific Reports*, 15(1):9174, 2025.
- [35] Aymen Qabel, Sofiane Ennadir, Giannis Nikolentzos, Johannes F Lutzeyer, Michail Chatzianastasis, Henrik Boström, and Michalis Vazirgiannis. Structure-aware antibiotic resistance classification using graph neural networks. In *NeurIPS 2022 AI for Science: Progress and Promises*, 2022.
- [36] Leslie Rice, Anna Bair, Huan Zhang, and J Zico Kolter. Robustness between the worst and average case. *Advances in Neural Information Processing Systems*, 34:27840–27851, 2021.
- [37] Lichao Sun, Yingdong Dou, Carl Yang, Kai Zhang, Ji Wang, Philip S. Yu, Lifang He, and Bo Li. Adversarial attack and defense on graph data: A survey. *IEEE Transactions on Knowledge and Data Engineering*, 2022.
- [38] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph Attention Networks. In *ICLR*, 2018.
- [39] Xingchen Wan, Henry Kenlay, Robin Ru, Arno Blaas, Michael A Osborne, and Xiaowen Dong. Adversarial attacks on graph classifiers via bayesian optimisation. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 6983–6996. Curran Associates, Inc., 2021.
- [40] Pengyun Wang, Junyu Luo, Yanxin Shen, Siyu Heng, and Xiao Luo. A comprehensive graph pooling benchmark: Effectiveness, robustness and generalizability. *CoRR*, 2024.
- [41] Per-Åke Wedin. Perturbation bounds in connection with singular value decomposition. *BIT Numerical Mathematics*, 12:99–111, 1972.
- [42] Huijun Wu, Chen Wang, Yuriy Tyshetskiy, Andrew Docherty, Kai Lu, and Liming Zhu. Adversarial examples for graph data: Deep insights into attack and defense. In *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19*, pages 4816–4823, 2019.
- [43] Kaidi Xu, Hongge Chen, Sijia Liu, Pin-Yu Chen, Tsui-Wei Weng, Mingyi Hong, and Xue Lin. Topology attack and defense for graph neural networks: An optimization perspective. *arXiv preprint arXiv:1906.04214*, 2019.
- [44] Keyulu Xu, Weihua Hu, Jure Leskovec, and Stefanie Jegelka. How Powerful are Graph Neural Networks? In *7th International Conference on Learning Representations*, 2019.
- [45] Muhan Zhang, Zhicheng Cui, Marion Neumann, and Yixin Chen. An end-to-end deep learning architecture for graph classification. In *Proceedings of the AAAI conference on artificial intelligence*, volume 32, 2018.
- [46] Xiang Zhang and Marinka Zitnik. Gnn-guard: Defending graph neural networks against adversarial attacks. *Advances in Neural Information Processing Systems*, 33:9263–9275, 2020.
- [47] Daniel Zügner, Oliver Borchert, Amir Akbarnejad, and Stephan Günnemann. Adversarial attacks on graph neural networks: Perturbations and their patterns. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 14(5):1–31, 2020.
- [48] Daniel Zügner and Stephan Günnemann. Adversarial attacks on graph neural networks via meta learning. In *7th International Conference on Learning Representations*, 2019.

- [49] Daniel Zügner, Amir Akbarnejad, and Stephan Günnemann. Adversarial attacks on neural networks for graph data. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM, 2018.
- [50] Daniel Zügner and Stephan Günnemann. Certifiable robustness and robust training for graph convolutional networks. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. ACM, jul 2019.

Supplementary Material: Enhancing Graph Classification Robustness with Singular Pooling

A Proof of Theorem 4.2

Theorem 4.2 Let $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ denote a graph-based function composed of L GCN layers, where the weight matrix of the ℓ -th layer is denoted by $W^{(\ell)}$. Under a feature-based adversarial attack with perturbation budget ϵ , the robustness of f in the sense of Definition 4.1 satisfies the following:

- If f is based on Max pooling, then it is (ϵ, γ) robust with:

$$\gamma = \sqrt{\min\{n, d_L\}} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \max_{u \in V} \hat{w}_u \epsilon.$$

- If f is based on Sum pooling, then it is (ϵ, γ) robust with:

$$\gamma = \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u \in V} \hat{w}_u \epsilon.$$

- If f is based on Average pooling, then it is (ϵ, γ) robust with:

$$\gamma = \frac{\epsilon}{n} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u \in V} \hat{w}_u.$$

Here, $\hat{w}_u$ denotes the sum of normalized walks of length $L-1$ originating from node u , and d_L is the output dimensionality of the final layer.

Proof. In this proof, we consider that f is a graph-function that is based on L layers of GCN. We recall that the GCN message-passing propagation is formulated for a node u as

$$h_u^{(\ell)} = \sigma^{(\ell)} \left(\sum_{v \in \mathcal{N}(u) \cup \{u\}} \frac{W^{(\ell)} h_v^{(\ell-1)}}{\sqrt{(1+d_u)(1+d_v)}} \right), \quad (4)$$

where $W^{(\ell)} \in \mathbb{R}^{d_{\ell-1} \times d_\ell}$ is the learnable weight matrix with d_ℓ being the embedding dimension of layer ℓ and $\sigma^{(\ell)}$ is the activation function of ℓ -th layer. We recall that $h^{(0)} = X \in \mathbb{R}^{n \times d}$ is set to the initial node features.

Similar to [1], we denote X as the original node features and denote by X' the perturbed adversarial features. We consider a node $u \in V$, we denote by h_u its representation in the clean graph and h'_u its representation in the attacked graph. We consider that the activation functions $(\sigma^{(\ell)})_{1 \leq \ell \leq L}$ are *nonexpansive* (1-Lipschitz continuous). From the work, we have the following result

$$\begin{aligned} \|h_u^{(L)} - h_{u'}^{(L)}\| &\leq \prod_{\ell=1}^L \|W^{(\ell)}\|_2 \left\| \sum_{v \in \mathcal{N}(u) \cup \{u\}} \sum_{j \in \mathcal{N}(v) \cup \{v\}} \cdots \right. \\ &\quad \left. \sum_{z \in \mathcal{N}(y) \cup \{y\}} \frac{X_u - X'_u}{\sqrt{(1+d_u)(1+d_w)(1+d_j) \cdots (1+d_y)\sqrt{(1+d_z)}}} \right\| \\ &\leq \prod_{\ell=1}^L \|W^{(\ell)}\| \hat{w}_u \epsilon, \end{aligned}$$

with $\hat{w}_u$ being the sum of normalized walks of length $(L-1)$ starting from node u .

The previous results give us an idea about the behavior of each node's representation when attacked. In the case of graph classification, an additional pooling operation is added, Specifically:

$$h_{\text{graph}}^{(L)} = \text{Pool} \left(\{h_u^{(L)}\}_{u \in V} \right).$$

Hence this proof's goal is to analyze the following quantity

$$\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\|.$$

Let's consider the **Sum pooling** operation, that can be written as

$$h_{\text{graph}}^{(L)} = \sum_{u \in V} h_u^{(L)}.$$

We have the following

$$\begin{aligned} \|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| &= \left\| \sum_{u \in V} (h_u^{(L)} - h'_u{}^{(L)}) \right\| \\ &\leq \sum_{u \in V} \|h_u^{(L)} - h'_u{}^{(L)}\| \quad (\text{by triangle inequality}) \\ &\leq \sum_{u \in V} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \|\hat{w}_u\| \epsilon \\ &= \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \epsilon \sum_{u \in V} \hat{w}_u. \end{aligned}$$

For the case of the **Average pooling** operation, that can be written as

$$h_{\text{graph}}^{(L)} = \frac{1}{|V|} \sum_{u \in V} h_u^{(L)}.$$

We have the following analysis

$$\begin{aligned} \|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| &= \left\| \frac{1}{|V|} \sum_{u \in V} (h_u^{(L)} - h'_u{}^{(L)}) \right\| \\ &\leq \frac{1}{|V|} \sum_{u \in V} \|h_u^{(L)} - h'_u{}^{(L)}\| \quad (\text{by triangle inequality}) \\ &\leq \frac{1}{|V|} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \epsilon \sum_{u \in V} \hat{w}_u. \end{aligned}$$

In the case of **Max pooling**: Let's denote $\Delta = H^{(L)} - H'^{(L)}$. Since the magnitude of a single entry of a row can never exceed the Euclidean norm of the whole row, which can be formulated as for any column j of our difference matrix that $|\Delta_{u,j}| \leq \|\Delta_{u,:}\|$. For Max pooling, we have the following

$$\begin{aligned} \|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\|_2 &\leq \sqrt{d_L} \max_u \|\Delta_{u,:}\|_2 \\ &\leq \epsilon \sqrt{d_L} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \max_{u \in V} \hat{w}_u. \end{aligned}$$

We additionally note that we have the following

$$\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\|^2 = \sum_{j=1}^{d_L} \max_u |\Delta_{u,j}|^2 \leq \sum_{u,j} |\Delta_{u,j}|^2 = \|\Delta\|_F^2. \quad (5)$$

And we have that

$$\|\Delta\|_F \leq \epsilon \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \left(\sum_u (\hat{w}_u^{(L)})^2 \right)^{1/2}.$$

From the two derived inequalities, we conclude that in the case of Max pooling we have

$$\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| \leq \epsilon \sqrt{\min\{n, d_L\}} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \max_{u \in V} \hat{w}_u.$$

By taking into account the expectancy (as shown in Definition 1), we get the desired results. $\square$

B Proof of Theorem 4.3

Theorem 4.3 Let $f: (\mathcal{A}, \mathcal{X}) \rightarrow \mathcal{Y}$ be composed of L GIN-layers (with its internal parameter $\zeta = 0$) and let $W^{(\ell)}$ denote the weight matrix of the ℓ -th MLP layer. We consider the input node feature space to be bounded, i.e. $\|X\|_2 < B$ for some $B \in \mathbb{R}$. Under a feature-based adversarial attack with perturbation budget ϵ , the model f is (ϵ, γ) -robust with respect to Definition 4.1, with:

- If f is based on Max pooling: $\gamma = \prod_{\ell=1}^L \|W^{(\ell)}\| (BL (\max_{u \in V} \deg(u)) + \epsilon)$.
- If f is based on Sum pooling: $\gamma = \prod_{\ell=1}^L \|W^{(\ell)}\| (2BL|E| + n\epsilon)$.
- If f is based on Average pooling: $\gamma = \sqrt{n, d_L} (\prod_{\ell=1}^L \|W^{(\ell)}\|) \left(\frac{2BL|E|}{n} + \epsilon \right)$.

Here, $|E|$ is the number of edges, n number of nodes and d_L the dimensionality of the final layer L .

Proof. In this proof, we consider that f is graph-based function that is based on L GIN-layers (with a parameter $\zeta = 0$, usually denoted as ϵ , but changed to avoid confusion with our attack budget). The GIN message-passing propagation process can be written for a node u as:

$$h_u^{(\ell+1)} = T^{(\ell+1)}((1 + \zeta)h_u^{(\ell)} + \sum_{v \in \mathcal{N}(u)} h_v^{(\ell)}),$$

with T denoting a Neural Networks (usually a MLP) and ζ denotes the parameter of the GIN. We recall that $h^{(0)} = X \in \mathbb{R}^{n \times d}$ is set to the initial node features.

Similar to the previous proof, we base our analysis on previous work [1], we denote X as the original node features and denote by X' the perturbed adversarial features. We consider a node $u \in V$, we denote by h_u its representation in the clean graph and h'_u its representation in the attacked graph. We recall that in our problem setup, we consider that the activation functions $(\sigma^{(\ell)})_{1 \leq \ell \leq L}$ are *nonexpensive* (1-Lipschitz continuous).

We use the same assumptions as the one considered in [1]. Specifically, we consider that the input feature space $\mathcal{H}_0$ is bounded, thus each hidden space $\mathcal{H}_i$ of the iterative process of message passing is bounded and let $B = \max_{\ell \leq L} B_\ell$ be its global maximum bound. We additionally consider that GIN-parameter $\zeta \approx 0$ (which is very frequent in the literature). For a node $u \in V$, we have the following result:

$$\|h_u^{(\ell+1)} - h'_{u'}^{(\ell+1)}\| \leq \prod_{\ell=1}^L \|W^{(\ell)}\| (BL \deg(u) + \epsilon).$$

From this perspective, let's consider the case of graph classification. We omit the definition of the different pooling as this was stated in the previous proof. We start by considering the **Sum pooling**

operation:

$$\begin{aligned}
\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| &= \left\| \sum_{u \in V} \left(h_u^{(L)} - h'_u{}^{(L)} \right) \right\| \\
&\leq \sum_{u \in V} \|h_u^{(L)} - h'_u{}^{(L)}\| \quad (\text{by the triangle inequality}) \\
&\leq \prod_{\ell=1}^L \|W^{(\ell)}\| \sum_{u \in V} (BL \deg(u) + \epsilon) \\
&= \prod_{\ell=1}^L \|W^{(\ell)}\| \left(BL \sum_{u \in V} \deg(u) + |V|\epsilon \right).
\end{aligned}$$

In the case of undirected graph, since $\sum_{u \in V} \deg(u) = 2|E|$, based on the previous result, we can write:

$$\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| \leq \prod_{\ell=1}^L \|W^{(\ell)}\| (2BL|E| + |V|\epsilon).$$

Similar for the case of **Average pooling**, we have:

$$\begin{aligned}
\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| &= \frac{1}{|V|} \left\| \sum_{u \in V} \left(h_u^{(L)} - h'_u{}^{(L)} \right) \right\| \\
&\leq \frac{1}{|V|} \sum_{u \in V} \|h_u^{(L)} - h'_u{}^{(L)}\| \\
&\leq \prod_{\ell=1}^L \|W^{(\ell)}\| \left(\frac{2BL|E|}{|V|} + \epsilon \right).
\end{aligned}$$

In the case of **Max pooling** operation, following the derivated Inequality 5, we can directly get:

$$\|h_{\text{graph}}^{(L)} - h'_{\text{graph}}^{(L)}\| \leq \sqrt{n, d_L} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \left(BL \left(\max_{u \in V} \deg(u) \right) + \epsilon \right).$$

By taking into account the expectancy (as shown in Definition 1), we get the desired results. $\square$

C Proof of Theorem 5.1

Theorem 5.1 Let $f: (\mathcal{G}, \mathcal{X}) \rightarrow \mathcal{Y}$ denote a graph-based function composed of L GCN layers and using our RS-Pool, where the weight matrix of the i -th layer is denoted by $W^{(i)}$. Under a feature-based adversarial attack with perturbation budget ϵ , f is (ϵ, γ) -robust in the sense of Definition 4.1 with:

$$\gamma = \frac{\tau\sqrt{2}\epsilon}{\sigma_1 - \sigma_2} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u=1}^n (\hat{w}_u)^2,$$

with $\hat{w}_u$ denoting the sum of normalized walks of length $(L-1)$ starting from node u , and $\sigma_1 \neq \sigma_2$ being the two dominant singular values.

Proof. In this proof, we consider that f is a graph-function that is based on L layers of GCN. Similar to previous proof, we denote by X' the perturbed node features, and $h'_u{}^{(\ell)}$ the corresponding node embedding of u at layer ℓ . Following [1], for each node u we have:

$$\|h_u^{(L)} - h'_u{}^{(L)}\| \leq \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \hat{w}_u \epsilon, \quad (6)$$

Thus, to relate $\|H^{(L)} - H'^{(L)}\|$ in operator norm

$$\|H^{(L)} - H'^{(L)}\|_F^2 = \sum_{u=1}^n \|h_u^{(L)} - h'_u{}^{(L)}\|^2 \leq \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right)^2 \epsilon^2 \sum_{u=1}^n (\hat{w}_u)^2.$$

Consequently, we have

$$\|H^{(L)} - H'^{(L)}\| \leq \|H^{(L)} - H'^{(L)}\|_F \leq \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \epsilon \sum_{u=1}^n (\hat{w}_u)^2. \quad (7)$$

We consider now our proposed pooling RS-Pool, which consists of considering the dominant right singular vector of $H^{(L)}$. In this perspective, let:

$$S = \text{span}\{v_1(H^{(L)})\} \text{ and } S' = \text{span}\{v_1(H'^{(L)})\} \text{ and accordingly: } \theta = \Theta(S, S') \in [0, \frac{\pi}{2}],$$

with θ being the principle angle between the two vectors. Note that we ensure to always take the smaller of the two possible angles between the two dominant vectors (by controlling the signs). Building on Wedin's sin- Θ theorem [41], we can write:

$$\sin \theta \leq \frac{\|H - H'\|}{\sigma_1(H) - \sigma_2(H)}.$$

From another perspective, for our considered subspaces S and S' , when considering the principal angle, we have the following:

$$\cos \theta = \max_{\substack{x \in S, y \in S' \\ \|x\|=\|y\|=1}} \langle x, y \rangle = |\langle v_1, v'_1 \rangle|.$$

We may assume $\langle v_1, v'_1 \rangle \geq 0$, since the sign of the dominant singular vector can be chosen by definition. From a different perspective, we can express as follows:

$$\begin{aligned} \|v_1 - v'_1\|^2 &= \langle v_1 - v'_1, v_1 - v'_1 \rangle = 2 - 2\langle v_1, v'_1 \rangle \\ &= 2 - 2\cos \theta \\ &= 4\sin^2(\theta/2). \end{aligned}$$

We additionally have the following (note that $\theta \in [0, \frac{\pi}{2}]$):

$$\sin \theta = 2 \sin \frac{\theta}{2} \cos \frac{\theta}{2} \Rightarrow \sin \frac{\theta}{2} = \frac{\sin \theta}{2 \cos \frac{\theta}{2}} \leq \frac{\sin \theta}{2 \cos \frac{\sqrt{2}}{2}} = \frac{\sin \theta}{\sqrt{2}}.$$

From the two computed elements, we find that:

$$\begin{aligned} \|v_1 - v'_1\|^2 &\leq 2 \sin(\theta/2) \\ &\leq 2 \frac{\sin \theta}{\sqrt{2}} \\ &\leq \sqrt{2} \sin \theta \\ &\leq \sqrt{2} \frac{\|H - H'\|}{\sigma_1(H) - \sigma_2(H)} \\ &\leq \frac{\sqrt{2}\epsilon}{\sigma_1 - \sigma_2} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u=1}^n (\hat{w}_u)^2. \end{aligned}$$

By also incorporating the constant $\tau \in \mathbb{R}_{>0}$, we obtain the final bound:

$$\gamma = \frac{\tau \sqrt{2}\epsilon}{\sigma_1 - \sigma_2} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u=1}^n (\hat{w}_u)^2.$$

On the special case of singular value degeneracy. In the previous analysis, we consider that the dominant singular value doesn't have a multiplicity higher than 1. We consider that this is the most dominant case in practice. Specifically, we note that RS-Pool is applied to the final node embedding matrix produced by the message-passing scheme. At this stage, having a dominant singular value with high multiplicity is highly unlikely, even in symmetric graphs, because the learned embeddings do not depend just on the graph structure but also on initial node features and the GNN parameters, which typically break such dominant symmetries. Moreover, GNNs are known for their inherent smoothing tendencies, which lead to rank-reduction in the feature matrix and result in a non-trivial spectral gap in practice.

Nonetheless, extending to the case where the dominant singular value is of multiplicity r shall result in an additional scalar $\sqrt{r}$ (rather than $\sqrt{2}$) in the derived upper-bound γ and the considered spectral gap is $\sigma_r - \sigma_{r+1}$.

□

Corollary 5.2 Let $f : (\mathcal{G}, \mathcal{X}) \rightarrow \mathcal{Y}$ be the graph-based function considered in Theorem 5.1. Then f is (ϵ, γ') -robust with:

$$\gamma' = \min\{\gamma, 2\tau\},$$

where γ is the bound derived in Theorem 5.1.

Proof. Let's consider the model $f : (\mathcal{G}, \mathcal{X}) \rightarrow \mathcal{Y}$ that uses our proposed RS-Pool.

From Theorem 5.1, we have seen that f is (ϵ, γ) -robust with:

$$\gamma = \frac{\tau\sqrt{2}\epsilon}{\sigma_1 - \sigma_2} \left(\prod_{\ell=1}^L \|W^{(\ell)}\| \right) \sum_{u=1}^n (\hat{w}_u)^2.$$

Using the triangular inequality, we have the following:

$$\begin{aligned} \|h_G - h_{G'}\| &= \|\text{RS-Pool}(H) - \text{RS-Pool}(H')\| \\ &= \|\tau v_1(H) - \tau v_1(H')\| \\ &= \tau \|v_1(H) - v_1(H')\| \\ &\leq \tau [\|v_1(H)\| + \|v_1(H')\|] \\ &\leq 2\tau. \end{aligned}$$

And therefore by combining the two provided bounds, we can write that f is (ϵ, γ') -robust with:

$$\gamma' = \min\{\gamma, 2\tau\}.$$

□

D Proof of Lemma 5.3

Lemma 5.3. RS-Pool is permutation invariant; that is, for any pair of isomorphic graphs G and G_Π with its corresponding permutation matrix Π , we have $f(A, X) = f(A_\Pi, X_\Pi)$.

Proof. Let $X \in \mathbb{R}^{n \times d}$ be a matrix, let's consider its SVD decomposition:

$$X = U\Sigma V^T.$$

Let's X' be a permuted version of the adjacency matrix X , and let $P \in \mathbb{R}^{n \times n}$ be the corresponding permutation matrix. Since by definition $P^T P = P P^T = I_n$, we can write the following:

$$(PX)^T(PX) = X^T P^T P X = X^T X.$$

From the previous equality, we observe that $X^T X$ and $(PX)^T(PX)$ are identical. Consequently, they share the same largest eigenvalue and corresponding eigenspace. Thus, any unit-norm dominant eigenvector of PX must lie in the one-dimensional subspace spanned by $v_1(X)$. By enforcing a consistent sign convention during pooling (e.g., requiring the first non-zero entry to be positive), we obtain:

$$\text{RS-Pool}(X') = \text{RS-Pool}(PX) = \tau v_1(PX) = \tau v_1(X) = \text{RS-Pool}(X).$$

As a result, we obtain permutation invariance in our pooling operation when X represents the output of the message-passing stage. $\square$

E Additional Results

E.1 Convergence Analysis

As outlined in Section 5, our proposed RS-Pool method relies on estimating the dominant singular vector of an input matrix using the iterative power method. This classical algorithm begins with a randomly initialized vector and performs a series of matrix-vector multiplications, progressively aligning the estimate with the dominant singular vector of the matrix. The number of iterations, denoted by K , directly influences the quality of this approximation and, by extension, the effectiveness of the pooling operation.

In this section, we empirically investigate how the choice of K impacts the accuracy of the estimated singular vector. To this end, we compare the vector obtained after K iterations of the power method to the dominant singular vector computed via the more precise full SVD implemented in the gesdd routine of the LAPACK package. We use the ℓ_2 -norm to measure the distance between the estimated and true singular vectors. This allows us to quantify the convergence behavior of the power method and assess how quickly a sufficiently accurate approximation can be achieved.

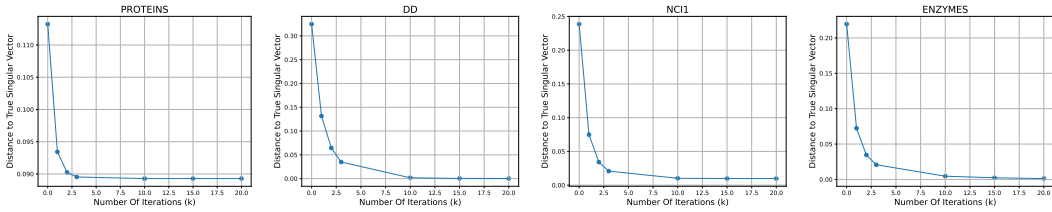


Figure 2: Effect of the number of iterations on the convergence of the iterative power method. The convergence is computed through the ℓ_2 distance between the estimated and true dominant singular vectors (obtained via full SVD).

Figure 2 summarizes the convergence behavior of our method across multiple datasets. We observe that in most cases, as few as 2 to 5 iterations are sufficient to obtain a reliable estimate of the dominant singular vector. This rapid convergence ensures that RS-Pool achieves stable performance in terms of both clean and adversarial accuracy, confirming its computational efficiency with a small number of iterations K . Moreover, the results indicate that the spectral gap $\sigma_1(X) - \sigma_2(X)$ is sufficiently large, which, according to Theorem 5.1, guarantees robustness to perturbations within a budget ϵ .

E.2 Time Analysis

In this section, we analyze the computational cost of our proposed RS-Pool method, with a particular focus on its time complexity. Building on the convergence analysis presented earlier, we examine how the number of power iterations K affects training time, clean accuracy, and adversarial robustness.

Table 3 reports results on the PROTEINS dataset. We observe that increasing the number of iterations yields only marginal improvements in clean accuracy. Similarly, the attack success rate decreases by less than 2% when increasing K from 1 to 10. In contrast, training time grows significantly due to the additional cost of repeated matrix-vector multiplications in the power iteration.

These results underscore the importance of selecting an appropriate number of iterations to balance model performance and computational efficiency. Taken together with the convergence properties

discussed earlier, those findings support that a small number of iterations (e.g., $K = 2$ to 5) is sufficient to achieve strong robustness with minimal overhead.

Table 3: Training time and Performance in terms of both clean and attack success rate for different number of iterations for the PROTEINS and NCI1 dataset for $\epsilon = 0.2$.

DATASET		$K = 1$	$K = 2$	$K = 5$	$K = 10$
PROTEINS	TRAINING TIME (IN S)	123.57	148.96	226.63	350.95
	CLEAN ACCURACY	73.36	73.95	73.51	73.06
	ATTACK SUCCESS RATE	21.04	21.25	20.50	19.89

We now focus on comparing the training time of the proposed RS-Pool against the benchmark pooling methods that were considered both in our theoretical study and our experimental results in Section 6.

Table 4: Mean training time (seconds) of the proposed RS-Pool in comparison to other pooling benchmarks on the different datasets.

Dataset	Sum	Avg	Max	SAG	TopK-P	Pan-P	Sort-P	RS-Pool
PROTEINS	56.95	56.50	60.78	81.27	71.32	68.98	71.59	148.96
D&D	1187.10	1176.93	1182.77	1206.33	1206.44	1210.92	1209.17	1293.93
NCI1	131.19	132.88	138.15	203.11	180.05	170.05	186.97	249.88
ENZYMES	20.42	20.57	21.54	30.76	27.54	26.34	28.28	46.30
IMDB-B	29.05	29.42	30.88	46.50	41.43	39.04	43.32	75.61
ER_MD	11.40	11.52	12.13	18.26	16.21	15.15	16.88	29.23
MSRC_9	6.62	6.71	7.06	9.93	8.96	8.39	9.31	18.43

E.3 Additional Results - GCN

Extending the analysis presented in Section 6, we continue our empirical evaluation of the upper bound γ across the remaining pooling methods. Consistent with earlier findings, RS-Pool exhibits a smaller distance between the pooled representations of clean and perturbed graphs compared to all other methods. This reduced discrepancy correlates with a lower attack success rate, highlighting RS-Pool’s improved adversarial robustness.

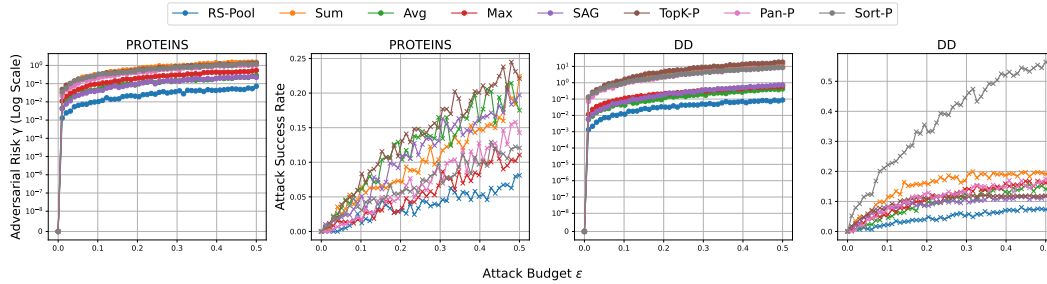


Figure 3: Extension of Figure 1: Empirical estimation of adversarial risk γ and corresponding attack success rate on PROTEINS ((a), (b)) and D&D ((c), (d)).

E.4 Additional Results - GIN

We further want to evaluate our proposed RS-Pool when considering a different backbone architecture, specifically Graph Isomorphism Networks (GIN). In this setup, we consider a model based on 2 GIN layers as the message-passing framework. Table 5 reports the mean and standard deviation of both clean and adversarial accuracy under PGD attacks with a perturbation budget of $\epsilon = 0.3$. Notably, we observe that in the case of GIN, increasing the number of power iteration steps used to approximate the dominant singular vector improves performance. As a result, we increase the number of iterations from $k = 2$ (used with GCN) to $k = 5$. Overall, RS-Pool consistently outperforms the baseline pooling methods, demonstrating similar robustness improvements to those observed

with GCN. These results support our claim that RS-Pool is model-agnostic and can be effectively integrated into different GNN architectures for graph classification.

Table 5: Clean and Attacked classification accuracy ($\pm$ standard deviation) of the RS-Pool and other pooling strategies applied on the GIN model when subject to the PGD adversarial attacks.

		Sum	Average	Max	SAG	TopK-P	PAN-P	Sort-P	RS-Pool
PROTEINS	Clean	75.6 $\pm$ 3.1	69.3 $\pm$ 1.5	75.8 $\pm$ 4.1	70.5 $\pm$ 1.9	73.5 $\pm$ 4.3	73.2 $\pm$ 2.6	73.2 $\pm$ 4.1	70.1 $\pm$ 1.6
	Attacked	22.6 $\pm$ 1.7	26.2 $\pm$ 2.9	41.3 $\pm$ 1.9	32.1 $\pm$ 5.9	26.5 $\pm$ 5.9	21.1 $\pm$ 3.6	35.8 $\pm$ 3.2	43.1 $\pm$ 3.9
D&D	Clean	73.5 $\pm$ 2.9	72.9 $\pm$ 3.1	73.5 $\pm$ 3.4	71.8 $\pm$ 0.9	74.4 $\pm$ 4.1	72.1 $\pm$ 4.9	73.8 $\pm$ 1.9	70.9 $\pm$ 3.6
	Attacked	12.9 $\pm$ 3.7	28.9 $\pm$ 5.6	5.6 $\pm$ 2.9	24.1 $\pm$ 4.9	23.9 $\pm$ 5.1	27.5 $\pm$ 3.6	32.9 $\pm$ 3.2	43.7 $\pm$ 5.1
ENZYMES	Clean	33.9 $\pm$ 4.9	27.8 $\pm$ 0.8	37.2 $\pm$ 4.3	27.8 $\pm$ 4.1	34.4 $\pm$ 0.8	30.0 $\pm$ 3.6	23.9 $\pm$ 2.8	27.8 $\pm$ 4.1
	Attacked	1.1 $\pm$ 0.8	0.6 $\pm$ 0.8	2.8 $\pm$ 2.1	2.5 $\pm$ 1.2	2.8 $\pm$ 0.8	5.6 $\pm$ 1.2	5.6 $\pm$ 0.8	6.1 $\pm$ 2.1
IMDB-B	Clean	64.3 $\pm$ 2.1	61.7 $\pm$ 4.7	52.3 $\pm$ 4.5	63.7 $\pm$ 4.1	61.0 $\pm$ 5.7	60.3 $\pm$ 4.7	66.0 $\pm$ 3.8	63.3 $\pm$ 4.7
	Attacked	43.7 $\pm$ 4.2	40.0 $\pm$ 5.1	42.0 $\pm$ 4.5	38.3 $\pm$ 4.3	34.3 $\pm$ 4.3	43.9 $\pm$ 2.1	35.0 $\pm$ 4.6	51.0 $\pm$ 3.8
ER_MD	Clean	68.2 $\pm$ 0.6	72.7 $\pm$ 2.8	71.2 $\pm$ 1.9	69.3 $\pm$ 4.6	67.8 $\pm$ 3.8	68.9 $\pm$ 2.8	61.8 $\pm$ 5.6	63.6 $\pm$ 2.4
	Attacked	34.1 $\pm$ 3.2	35.6 $\pm$ 2.9	34.1 $\pm$ 1.4	36.7 $\pm$ 3.2	36.7 $\pm$ 3.2	35.6 $\pm$ 1.4	40.8 $\pm$ 2.3	60.1 $\pm$ 2.8
MSRC_9	Clean	93.9 $\pm$ 2.1	89.4 $\pm$ 2.1	93.2 $\pm$ 4.9	91.7 $\pm$ 1.1	93.2 $\pm$ 1.9	92.4 $\pm$ 1.1	91.7 $\pm$ 2.8	91.7 $\pm$ 4.6
	Attacked	65.9 $\pm$ 7.4	72.7 $\pm$ 3.7	60.6 $\pm$ 2.1	67.5 $\pm$ 3.8	69.6 $\pm$ 5.2	54.6 $\pm$ 3.7	65.2 $\pm$ 7.7	83.3 $\pm$ 4.2

E.5 Additional Results - Adversarial Defenses

Since RS-Pool is model-agnostic, we investigate its effectiveness when combined with existing adversarial defense strategies. We consider four representative defense categories. First, a *post-processing* defense via adversarial training, where adversarial examples are generated and used to retrain the model to improve robustness. We additionally consider Randomized smoothing, which consists of generating a number of points within the input and then take the majority vote as a classification. We afterwards consider a set of *architecture-level* defense using both R-GCN, which modifies the message-passing mechanism to improve stability under perturbations, and SoftMedian, which introduces a robust aggregation scheme that down-weights neighbors based on their distance to the dimension-wise median, mitigating the impact of outliers or malicious inputs.

Table 6 provides the mean clean and attacked accuracy for the different considered pooling benchmarks when subject to the PDG adversarial attack. We observe similar insights as in the case of the GCN and GIN, where our proposed RS-Pool out-perform all the other poolings in terms of attacked accuracy showcasing therefore its ability to work with different underlying backbone GNNs.

E.6 Bit-Flip Attacks (BFAs)

Recently, Bit-Flip attacks have emerged as a new paradigm for attacking and consequently evaluating the robustness of deep learning models at the hardware and memory level. Unlike conventional adversarial attacks that perturb inputs or gradients, bit-flip attacks directly manipulate the binary representation of a model’s parameters in memory. By flipping only a few critical bits, an adversary can drastically alter the network’s behavior, causing severe accuracy degradation or targeted misclassification. While such attacks have been extensively studied in the context of convolutional and feedforward architectures, their implications for Graph Neural Networks (GNNs) are only beginning to be understood. Recent work [26, 25] have shown that even minimal bit-level corruption can disrupt the delicate balance of the message-passing propagation mechanism.

In this perspective, and to further demonstrate the robustness and validity of our proposed RS-Pool, we extend our evaluation to include this line of attack. Specifically, we adapt the *Injectivity Bit-Flip Attack* (IBFA) by removing its quantization step, allowing it to operate directly on our floating-point model. We employ the same dataset and follow an identical experimental setup, using the GIN architecture and setting the attack parameter $k = 5$, consistent with the configuration used by the original authors (on the OGB-molHIV molecular dataset). The results of this analysis, summarized in Table 7, indicate that indeed RS-Pool consistently outperforms all other pooling baselines, further validating its resilience to bit-level adversarial perturbations.

Table 6: Clean and Attacked classification accuracy ($\pm$ standard deviation) of the RS-Pool and other pooling strategies applied on different graph adversarial methods when subject to the PGD adversarial attacks.

			Sum	Average	Max	SAG	TopK-P	PAN-P	Sort-P	RS-Pool
PROTEINS	Adv	Clean	72.4 $\pm$ 4.0	70.2 $\pm$ 2.3	73.2 $\pm$ 3.7	71.1 $\pm$ 3.7	73.7 $\pm$ 2.9	73.5 $\pm$ 2.2	73.8 $\pm$ 2.3	72.1 $\pm$ 3.1
	Train.	Attacked	45.8 $\pm$ 5.1	31.3 $\pm$ 2.2	18.5 $\pm$ 3.7	49.1 $\pm$ 3.6	46.1 $\pm$ 4.8	32.7 $\pm$ 3.1	51.5 $\pm$ 3.5	54.8 $\pm$ 4.1
	SoftMedian	Clean	72.6 $\pm$ 4.7	71.2 $\pm$ 1.4	73.2 $\pm$ 4.6	72.0 $\pm$ 0.9	73.2 $\pm$ 0.8	73.0 $\pm$ 5.9	73.7 $\pm$ 0.7	74.3 $\pm$ 1.3
		Attacked	53.9 $\pm$ 4.9	42.8 $\pm$ 6.7	41.5 $\pm$ 5.1	36.3 $\pm$ 2.9	52.4 $\pm$ 1.2	41.7 $\pm$ 2.1	58.3 $\pm$ 1.8	63.2 $\pm$ 2.8
	RGCN	Clean	72.6 $\pm$ 2.3	61.0 $\pm$ 1.6	57.6 $\pm$ 3.4	66.7 $\pm$ 2.2	70.2 $\pm$ 3.4	70.5 $\pm$ 1.9	70.8 $\pm$ 0.4	71.7 $\pm$ 1.1
		Attacked	65.4 $\pm$ 0.8	58.6 $\pm$ 1.7	39.8 $\pm$ 4.2	59.2 $\pm$ 1.1	63.1 $\pm$ 3.4	61.3 $\pm$ 4.7	63.6 $\pm$ 2.3	67.3 $\pm$ 1.8
	Randomized	Clean	75.4 $\pm$ 3.6	71.7 $\pm$ 3.3	72.6 $\pm$ 2.2	70.8 $\pm$ 2.3	74.1 $\pm$ 3.6	75.6 $\pm$ 3.1	74.4 $\pm$ 4.1	73.5 $\pm$ 3.4
	Smoothing	Attacked	56.2 $\pm$ 1.5	49.7 $\pm$ 3.7	36.3 $\pm$ 3.6	57.2 $\pm$ 4.4	56.1 $\pm$ 3.2	53.6 $\pm$ 0.7	52.9 $\pm$ 4.3	59.3 $\pm$ 4.8
	Adv	Clean	76.7 $\pm$ 3.4	68.9 $\pm$ 4.1	70.4 $\pm$ 4.3	69.3 $\pm$ 3.6	75.1 $\pm$ 4.3	74.9 $\pm$ 2.6	76.0 $\pm$ 2.5	74.7 $\pm$ 3.4
	Train.	Attacked	7.3 $\pm$ 2.7	18.6 $\pm$ 6.3	6.7 $\pm$ 2.2	19.6 $\pm$ 5.1	8.7 $\pm$ 1.7	7.1 $\pm$ 2.6	7.3 $\pm$ 1.5	31.9 $\pm$ 4.9
D&D	SoftMedian	Clean	74.1 $\pm$ 3.4	68.9 $\pm$ 4.7	69.0 $\pm$ 5.1	68.4 $\pm$ 5.1	75.7 $\pm$ 5.8	77.3 $\pm$ 1.6	74.9 $\pm$ 2.9	72.9 $\pm$ 3.4
		Attacked	14.9 $\pm$ 3.8	51.2 $\pm$ 4.3	8.5 $\pm$ 3.8	46.1 $\pm$ 5.6	10.9 $\pm$ 4.3	29.1 $\pm$ 6.4	10.4 $\pm$ 4.4	54.3 $\pm$ 3.8
	RGCN	Clean	74.9 $\pm$ 1.3	63.8 $\pm$ 2.4	66.4 $\pm$ 5.1	70.6 $\pm$ 3.4	75.4 $\pm$ 2.1	76.8 $\pm$ 1.8	73.1 $\pm$ 2.4	71.2 $\pm$ 2.8
		Attacked	37.8 $\pm$ 2.1	22.8 $\pm$ 3.6	13.4 $\pm$ 2.9	38.7 $\pm$ 4.8	44.5 $\pm$ 2.4	58.7 $\pm$ 3.4	53.7 $\pm$ 3.6	61.4 $\pm$ 3.2
	Randomized	Clean	74.9 $\pm$ 4.3	72.1 $\pm$ 2.8	42.5 $\pm$ 2.7	69.3 $\pm$ 1.7	74.8 $\pm$ 1.5	76.4 $\pm$ 1.5	76.1 $\pm$ 1.9	72.1 $\pm$ 3.8
	Smoothing	Attacked	39.9 $\pm$ 2.6	43.1 $\pm$ 3.5	22.8 $\pm$ 4.5	34.1 $\pm$ 4.1	45.1 $\pm$ 4.8	41.1 $\pm$ 3.5	42.2 $\pm$ 1.7	49.7 $\pm$ 4.6
ER_MD	Adv	Clean	63.7 $\pm$ 2.3	63.3 $\pm$ 3.4	63.3 $\pm$ 4.6	62.2 $\pm$ 4.5	63.3 $\pm$ 4.1	68.1 $\pm$ 0.9	65.5 $\pm$ 4.8	63.3 $\pm$ 2.8
	Train.	Attacked	42.6 $\pm$ 4.1	45.7 $\pm$ 2.9	44.2 $\pm$ 2.6	48.6 $\pm$ 1.4	42.3 $\pm$ 4.9	35.4 $\pm$ 4.2	34.9 $\pm$ 3.3	53.9 $\pm$ 2.6
	SoftMedian	Clean	70.0 $\pm$ 1.9	62.9 $\pm$ 6.9	62.5 $\pm$ 4.5	62.5 $\pm$ 4.0	65.1 $\pm$ 4.8	66.7 $\pm$ 6.8	65.5 $\pm$ 6.1	62.5 $\pm$ 1.9
		Attacked	36.3 $\pm$ 5.5	42.3 $\pm$ 5.0	49.1 $\pm$ 5.0	43.4 $\pm$ 4.5	42.3 $\pm$ 5.0	26.6 $\pm$ 3.4	34.5 $\pm$ 1.0	53.6 $\pm$ 3.4
	RGCN	Clean	58.0 $\pm$ 5.5	60.1 $\pm$ 4.3	58.4 $\pm$ 4.7	62.9 $\pm$ 5.7	63.6 $\pm$ 4.6	68.5 $\pm$ 6.4	66.7 $\pm$ 3.7	59.9 $\pm$ 3.7
		Attacked	39.2 $\pm$ 1.2	47.5 $\pm$ 2.1	44.5 $\pm$ 3.9	49.3 $\pm$ 5.1	28.0 $\pm$ 3.8	38.6 $\pm$ 1.4	46.8 $\pm$ 1.4	51.7 $\pm$ 0.9
	Randomized	Clean	69.2 $\pm$ 1.0	61.4 $\pm$ 3.4	61.4 $\pm$ 4.4	61.0 $\pm$ 3.5	62.9 $\pm$ 4.2	67.8 $\pm$ 2.3	63.6 $\pm$ 5.4	61.7 $\pm$ 4.1
	Smoothing	Attacked	40.4 $\pm$ 3.4	40.4 $\pm$ 6.4	46.4 $\pm$ 1.1	48.3 $\pm$ 3.4	40.1 $\pm$ 4.7	43.4 $\pm$ 5.1	43.8 $\pm$ 4.4	54.7 $\pm$ 4.4

Table 7: Attacked classification accuracy ($\pm$ standard deviation) of the RS-Pool and other pooling strategies on different graph classification dataset when subject to Bit-Flip Attacks.

Dataset	Sum	Average	Max	SAG	TopK-P	PAN-P	Sort-P	RS-Pool
NCII	64.7 $\pm$ 1.9	62.2 $\pm$ 1.9	65.9 $\pm$ 1.3	65.2 $\pm$ 1.3	61.2 $\pm$ 1.7	59.2 $\pm$ 1.7	63.5 $\pm$ 1.7	66.0 $\pm$ 1.6
DD	59.8 $\pm$ 4.1	61.8 $\pm$ 2.7	63.6 $\pm$ 2.6	63.1 $\pm$ 1.4	58.4 $\pm$ 2.1	56.2 $\pm$ 1.5	56.8 $\pm$ 2.5	68.3 $\pm$ 1.0
ER_MD	55.2 $\pm$ 2.5	52.3 $\pm$ 1.8	51.1 $\pm$ 1.8	52.4 $\pm$ 2.0	52.9 $\pm$ 2.7	52.2 $\pm$ 1.1	51.9 $\pm$ 2.7	55.5 $\pm$ 2.5
PROTEINS	57.0 $\pm$ 5.8	64.2 $\pm$ 1.5	64.8 $\pm$ 2.7	63.7 $\pm$ 0.9	61.2 $\pm$ 5.0	60.0 $\pm$ 4.9	66.7 $\pm$ 4.7	69.6 $\pm$ 1.1

E.7 On the Tightness of Upper-Bounds

In this section, we assess the tightness of the theoretical upper bound γ established in Theorem 4.2. To this end, we compare the theoretical bound with empirical estimates computed as the norm difference between the pooled representations of clean and perturbed graphs for different values of ϵ . The results are shown in Figure 4. Across all pooling methods, we observe a clear correlation between the empirical measurements and the theoretical bound, providing empirical support for the validity of our theoretical analysis.

F Experimental Details

F.1 Datasets

For our experiments, we use standard graph classification datasets from the TUDataset benchmark [32], selecting a diverse subset to represent different domains. Specifically, we include datasets from bioinformatics and chemoinformatics (PROTEINS, ENZYMES, D&D), small molecule prediction (NCII, ER_MD), and image-based graphs (MSRC_9). Model evaluation follows the standardized protocol of [14], using 10-fold cross-validation. When public folds are available, we use them directly;

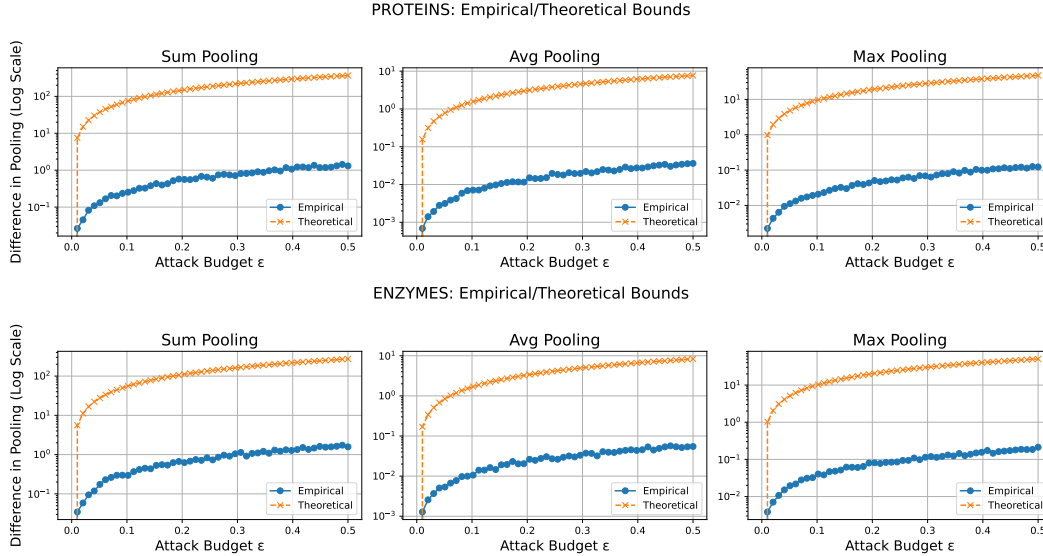


Figure 4: Analysis of the Upper-Bounds provided in Theorem 4.2 and their empirical counter-part for different values of ϵ on the PROTEINS (Upper) and ENZYMES (Lower) Dataset.

otherwise, we generate new folds following the same procedure. Dataset statistics and characteristics are summarized in Table 8.

Table 8: Statistics of the graph classification datasets used in our experiments.

DATASET	#GRAPHS	#NODES	#EDGES	#CLASSES
PROTEINS	1113	39.06	72.82	2
DD	1178	284.32	715.66	2
NCI1	4110	29.87	32.30	2
ENZYMES	600	32.63	62.14	6
IMDB-B	1000	19.77	96.53	2
REDDIT-B	2000	429.63	497.75	2
ER_MD	446	21.33	234.85	2
MSRC_9	221	40.58	97.94	8

F.2 Experimental Setup

In all of the experiments, we use a 2-Layers GCN model. Specifically, the model uses a 2-layer convolutional architecture (consisting of two iterations of message passing and updating) stacked with a Multi-Layer Perception (MLP) as a readout. The only change in the architecture was on the pooling level since the intent was to compare the different benchmarks in an iso-architectural setting, to ensure a fair evaluation of their robustness.

All models are trained using the Adam optimizer [23] with a learning rate of 1×10^{-3} for 100 epochs. We set the hidden feature dimension to 32 and we used ReLU as our activation function for all the models. For all the experiments, we used the same initialization distribution and the same number of training epochs to ensure fairness following insights from previous work[12]. Additionally, to account for variability due to random initialization, each experiment is repeated 10 times, and we report the mean and standard deviation of the results. The experiments have been run on a NVIDIA A40 GPU and we estimate the total number of hours of computing to be around 200 hours.

For our proposed RS-Pool, we set the number of iterations in the estimation algorithm to $K = 2$ in the case of GCN while we set it to $K = 5$ in the case of GIN. In addition, the temperature parameter τ was tuned separately for each dataset to optimize the trade-off between clean accuracy and adversarial robustness, as detailed in Section 6 (Figures 1 (e)).

F.3 Implementation Details

Our implementation is provided in the supplementary materials and will be made publicly available after the review period. The code is developed using PyTorch [33], with a dense implementation of GNNs, which is required for executing the considered adversarial attacks. For all baseline pooling methods, we adapt the official implementations from *PyTorch Geometric* (PyG) [15], released under the MIT license.

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: The main claims presented in the abstract and introduction are supported by the theoretical results in Section 4 and further validated through empirical evidence in Section 6.

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: We discuss the main limitations of our method in Section 7.

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: All theoretical claims are supported by detailed proofs in Appendices A, B, C, and D. Additionally, the problem setup and underlying assumptions are clearly defined in Section 3 and are explicitly referenced in each theorem and proof.

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: In addition to presenting a detailed experimental setup in Appendix F, we include the source code for reproducing our results in the supplementary materials.

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: All experiments were conducted using publicly available datasets. When available, we used the available public train/val/test folds, and we provide all of those. To ensure reproducibility, we provide the source code in the supplementary materials.

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: Appendix F provides details of the experimental setup, including details about the train/val/test used folds and the values of all hyperparameters.

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 report the mean and the standard deviation of all results in the respective tables, based on repeated experiments using 10 different random seeds. Further details are provided in Appendix F.

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: We provide information about compute resources for our experiments in Appendix F.

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: Our work adheres to the NeurIPS Code of Ethics. We use public datasets and publicly available models and report on the limitations of our work.

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: [NA]

Justification: Our work has no further societal impacts apart from known impacts of graph-based models. We are not aware of any applications of our approach that would result in negative societal impacts.

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: Our paper does not release any data or models that have a high risk for misuse.

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: All datasets, models and repositories were cited appropriately.

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: [NA]

Justification: Our paper does not release new assets, only rely on public datasets and models.

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: Justification: Our work did not involve crowdsourcing nor research with human subjects. All experiments are performed on publicly available datasets.

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: Our paper does not involve crowdsourcing nor research with human subjects, all experiments are performed on publicly available datasets.

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.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: No LLM was used in this work for the core methods.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.