
Effects of Dropout on Performance in Long-range Graph Learning Tasks

Jasraj Singh^{1*} Keyue Jiang^{2†} Brooks Paige² Laura Toni²
¹Nanyang Technological University ²University College London

Abstract

Message Passing Neural Networks (MPNNs) are a class of Graph Neural Networks (GNNs) that propagate information across the graph via local neighborhoods. The scheme gives rise to two key challenges: *over-smoothing* and *over-squashing*. While several Dropout-style algorithms, such as DropEdge and DropMessage, have successfully addressed over-smoothing, their impact on over-squashing remains largely unexplored. This represents a critical gap in the literature, as failure to mitigate over-squashing would make these methods unsuitable for long-range tasks – the intended use case of deep MPNNs. In this work, we study the aforementioned algorithms, and closely related edge-dropping algorithms – DropNode, DropAgg and DropGNN – in the context of over-squashing. We present theoretical results showing that DropEdge-variants reduce sensitivity between distant nodes, limiting their suitability for long-range tasks. To address this, we introduce DropSens, a sensitivity-aware variant of DropEdge, which is developed following the message-passing scheme of GCN. DropSens explicitly controls the proportion of information lost due to edge-dropping, thereby increasing sensitivity to distant nodes despite dropping the same number of edges. Our experiments on long-range synthetic and real-world datasets confirm the predicted limitations of existing edge-dropping and feature-dropping methods. Moreover, DropSens with GCN consistently outperforms graph rewiring techniques designed to mitigate over-squashing, suggesting that simple, targeted modifications can substantially improve a model’s ability to capture long-range interactions. Our conclusions highlight the need to re-evaluate and re-design existing methods for training deep GNNs, with a renewed focus on modelling long-range interactions. The code for reproducing the results in our this work is available at <https://github.com/ignasa007/Dropout-Effects-GNNs>.

1 Introduction

Graph neural networks (GNNs) [51, 73] are powerful neural models developed for modelling graph-structured data, and have found applications in several real-world scenarios [29, 61, 84, 90–94, 96, 101]. A popular class of GNNs, called *message-passing neural networks* (MPNNs) [32], recursively process neighborhood information using message-passing layers. These layers are stacked to allow each node to aggregate information from increasingly larger neighborhoods, akin to how *convolutional neural networks* (CNNs) learn hierarchical features for images [48]. However, unlike in image-based deep learning, where *ultra-deep* CNN architectures have led to performance breakthroughs [38, 80], shallow GNNs often outperform deeper models on many graph learning

*Part of the work done as a master’s student at UCL. †Primary supervisor. Correspondence to Jasraj Singh <jasraj.singh00150@gmail.com> and Keyue Jiang <keyue.jiang.18@ucl.ac.uk>.

tasks [99]. This is because deep GNNs suffer from unique issues like *over-smoothing* [65] and *over-squashing* [4], which makes training them notoriously difficult.

Over-smoothing refers to the problem of node representations becoming *too similar* as they are recursively processed. This is undesirable since it limits the GNN from effectively utilizing the information in the input features. The problem has garnered significant attention from the research community, resulting in a suite of algorithms designed to address it [72] (see Appendix A.1 for an overview of representative methods). Amongst these methods are a collection of random edge-dropping algorithms, including DropEdge [70], DropNode [23], DropAgg [43] and DropGNN [66] – which we will collectively refer to as *DropEdge-variants* – which act as *message-passing reducers*. In addition, we have DropMessage [21], which performs Dropout [79] on the message matrices, instead of the feature matrices; we will collectively refer to these two methods as *Dropout-variants* since they are applied along the feature dimensions.

The other issue specific to GNNs is over-squashing. In certain graph structures, neighborhood size grows exponentially with distance from the source [12], causing information to be lost as it passes through graph bottlenecks [4]. This limits MPNNs’ ability to enable communication between distant nodes, which is crucial for good performance on long-range tasks [52]. To alleviate over-squashing, several graph-rewiring techniques have been proposed, which aim to improve graph connectivity by adding edges in a strategic manner [4, 8, 16, 44, 64] (see Appendix A.3 for an overview of representative methods).² In contrast, the DropEdge-variants only remove edges, which should, in principle, amplify over-squashing levels. The same can be intuitively argued about Dropout-variants.

Empirical evidence in support of methods designed for training deep GNNs has been majorly collected on short-range tasks (see Appendix A.2 for a detailed discussion). That is, it simply suggests that *these methods prevent loss of local information, but it remains inconclusive if they facilitate capturing long-range interactions* (LRIs). Of course, on long-range tasks, deep GNNs are useless if they cannot capture LRIs. This is especially a concern for DropEdge-variants since evidence suggests that alleviating over-smoothing with graph rewiring could exacerbate over-squashing [34, 64].

Contributions. In this work, we precisely characterize the effects of random edge-dropping algorithms on over-squashing in MPNNs. By explicitly computing the expected *sensitivity* of the node representations to the node features [81] (inversely related to over-squashing) in a linear Graph Convolutional Network (GCN) [47], we show that these methods provably reduce the *effective receptive field* of the model. Precisely speaking, the rate at which sensitivity between nodes decays is *exponential* w.r.t. the distance between them. We also extend the existing theoretical results on sensitivity in nonlinear MPNNs [8, 18, 86] to the random edge-dropping setting, concluding that these algorithms exacerbate the over-squashing problem. We use our analysis of GCNs to design a sensitivity-aware DropEdge-variant, named *DropSens*, that enjoys the representational expressivity of DropEdge without suffering from over-squashing, thereby demonstrating how algorithms can be readily adapted for long-range tasks.

We evaluated the DropEdge- and Dropout-variants on long-range datasets using GCN, Graph Isomorphism Network (GIN) [87] and Graph Attention Network (GAT) [83] architectures. Specifically, we follow the setup in [33] with the SyntheticZINC dataset, in [81] with real-world homophilic (corresponding to short-range tasks) and heterophilic (long-range tasks) node classification datasets, and in [8, 44] with graph classification datasets. Our results indicate that while the random dropping methods improve model performance in short-range tasks, they are often ineffective, and sometimes even detrimental, to long-range task performance. Finally, we present results for DropSens, which outperforms state-of-the-art graph rewiring methods aimed at addressing over-squashing at node classification and graph-classification tasks. These findings point to the importance of re-evaluating the methods used to train deep GNNs, especially in terms of how well they capture LRIs.

2 Background

Consider a directed graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, with $\mathcal{V} = [N] := \{1, \dots, N\}$ denoting the node set and $\mathcal{E} \subset \mathcal{V} \times \mathcal{V}$ the edge set; $(j \rightarrow i) \in \mathcal{E}$ if there’s an edge from node j to node i . Let $\mathbf{A} \in \{0, 1\}^{N \times N}$ denote its adjacency matrix, such that $\mathbf{A}_{ij} = 1$ if and only if $(j \rightarrow i) \in \mathcal{E}$, and let $\mathbf{D} := \text{diag}(\mathbf{A}\mathbf{1}_N)$ denote the in-degree matrix. The geodesic distance, $d_{\mathcal{G}}(j, i)$, from node j to node i is the length

²Sometimes, along with removal of some edges to preserve statistical properties of the original topology.

of the shortest path starting at node j and ending at node i . Accordingly, the ℓ -hop neighborhood of a node i can be defined as the set of nodes that can reach it in exactly $\ell \in \mathbb{N}_0$ steps, $\mathbb{S}^{(\ell)}(i) = \{j \in \mathcal{V} : d_{\mathcal{G}}(j, i) = \ell\}$.

2.1 Graph Neural Networks

Graph Neural Networks (GNNs) operate on inputs of the form $(\mathcal{G}, \mathbf{X})$, where $\mathcal{G}$ encodes the graph topology and $\mathbf{X} \in \mathbb{R}^{N \times H^{(0)}}$ collects the node features.³ Message-Passing Neural Networks (MPNNs) [32] are a special class of GNNs which recursively aggregate information from the 1-hop neighborhood of each node using *message-passing layers*. An L-layer MPNN is given as

$$\begin{aligned} \mathbf{z}_i^{(\ell)} &= \text{Upd}^{(\ell)} \left(\mathbf{z}_i^{(\ell-1)}, \text{Agg}^{(\ell)} \left(\mathbf{z}_i^{(\ell-1)}, \left\{ \mathbf{z}_j^{(\ell-1)} : j \in \mathbb{S}^{(1)}(i) \right\} \right) \right), \quad \forall \ell \in [L] \\ \text{MPNN}_{\theta}(\mathcal{G}, \mathbf{X}) &= \text{Out} \left(\left\{ \mathbf{z}_i^{(L)} : i \in \mathcal{V} \right\} \right) \end{aligned} \quad (2.1)$$

where $\mathbf{Z}^{(0)} = \mathbf{X}$, $\text{Agg}^{(\ell)}$ denotes the *aggregation functions*, $\text{Upd}^{(\ell)}$ the *update functions*, and Out the *readout function*. Since $\mathbf{z}_i^{(L)}$ is a function of the input features of nodes at most L-hops away from it, its *receptive field* is given by $\mathbb{R}^{(L)}(i) := \{j \in \mathcal{V} : d_{\mathcal{G}}(j, i) \leq L\}$.

For example, a GCN [47] updates node representations as the weighted sum of its neighbors' representations:

$$\mathbf{Z}^{(\ell)} = \sigma \left(\hat{\mathbf{A}} \mathbf{Z}^{(\ell-1)} \mathbf{W}^{(\ell)} \right) \quad (2.2)$$

where σ is a point-wise nonlinearity, e.g. ReLU, the propagation matrix, $\hat{\mathbf{A}}$, is a *graph shift operator*, i.e. $\hat{\mathbf{A}}_{ij} \neq 0$ if and only if $(j \rightarrow i) \in \mathcal{E}$ or $i = j$, and $\mathbf{W}^{(\ell)} \in \mathbb{R}^{H^{(\ell-1)} \times H^{(\ell)}}$ is a weight matrix. The original choice for $\hat{\mathbf{A}}$ was the symmetrically normalized adjacency matrix $\hat{\mathbf{A}}^{\text{sym}} := \tilde{\mathbf{D}}^{-1/2} \tilde{\mathbf{A}} \tilde{\mathbf{D}}^{-1/2}$ [47], where $\tilde{\mathbf{A}} = \mathbf{A} + \mathbf{I}_N$ and $\tilde{\mathbf{D}} = \text{diag}(\tilde{\mathbf{A}} \mathbf{1}_N)$. However, several influential works have also used the asymmetrically normalized adjacency, $\hat{\mathbf{A}}^{\text{asym}} := \tilde{\mathbf{D}}^{-1} \tilde{\mathbf{A}}$ [36, 50, 74].

2.2 DropEdge-variants

DropEdge [70] is a random data augmentation technique that works by sampling a subgraph of the input graph in each layer, followed by the addition of self-loops, and uses that for message passing. Several variants of DropEdge have also been proposed, forming a family of random edge-dropping algorithms for tackling the over-smoothing problem. For example, DropNode [23] independently samples nodes and sets their features to $\mathbf{0}$, followed by rescaling to make the feature matrix unbiased. This is equivalent to setting the corresponding columns of the propagation matrix to $\mathbf{0}$. In a similar vein, DropAgg [43] samples nodes that don't aggregate messages from their neighbors. This is equivalent to dropping the corresponding rows of the adjacency matrix. Combining these two approaches, DropGNN [66] samples nodes which neither propagate nor aggregate messages in a given layer. These algorithms alleviate over-smoothing by reducing the number of messages being propagated in the graph, thereby slowing down the convergence of node representations.

2.3 Dropout-variants

Dropout is a stochastic regularization technique which reduces over-fitting by randomly dropping features before each layer. It has been successful with various architectures, like CNNs [79] and transformers [82], and has also found applications in GNN training. DropMessage [21] is a variant of Dropout designed specifically for message-passing schemes – it acts directly on the messages over each edge, instead of the node representations. This reduces the induced variance in the messages compared to Dropout, DropEdge and DropNode, while at the same time making the method more effective at alleviating over-smoothing and enabling the training of deep GNNs.

³To keep things simple, we will ignore edge features.

2.4 Over-squashing

Over-squashing refers to the problem of information from exponentially growing neighborhoods [11] being squashed into finite-sized node representations [4]. [81] formally characterized over-squashing in terms of the Jacobian of the node-level representations w.r.t. the input features: $\|\partial \mathbf{z}_i^{(L)} / \partial \mathbf{x}_j\|_1$. Accordingly, over-squashing can be understood as low sensitivity between distant nodes, i.e. small perturbations in a node's features don't effect other distant nodes' representations.

See Appendix A for an extensive discussion of related works addressing the problems of over-smoothing and over-squashing, and a unified treatment of the two.

3 Sensitivity Analysis

In this section, we perform a theoretical analysis of the expectation – w.r.t. random edge masks – of sensitivity of node representations. This will allow us to predict how DropEdge-variants affect communication between nodes at various distances, which is relevant for predicting their suitability towards learning LRIs.

Here, we present our analysis for linear GCNs, and treat more general nonlinear MPNN architectures in Appendix C.1. In this model, the final node representations can be summarised as

$$\mathbf{Z}^{(L)} = \left(\prod_{\ell=1}^L \hat{\mathbf{A}}^{(\ell)} \right) \mathbf{X} \mathbf{W} \in \mathbb{R}^{N \times H^{(L)}} \quad (3.1)$$

where $\mathbf{W} := \prod_{\ell=1}^L \mathbf{W}^{(\ell)} \in \mathbb{R}^{H^{(0)} \times H^{(L)}}$. Using the i.i.d. assumption on the distribution of edge masks in each layer, $\{\mathbf{M}^{(\ell)}\}_{\ell=1}^L$, the expected sensitivity of node i to node j can be shown to be

$$\mathbb{E}_{\mathbf{M}^{(1)}, \dots, \mathbf{M}^{(L)}} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] = \left(\mathbb{E} [\hat{\mathbf{A}}] \right)_{ij}^L \|\mathbf{W}\|_1 \quad (3.2)$$

To keep things simple, we will ignore the effect of DropEdge-variants on the optimization trajectory. Accordingly, it is sufficient to study $\mathbb{E}[\hat{\mathbf{A}}]$ in order to predict their effect on over-squashing. To maintain analytical tractability, we assume the use of an asymmetrically normalized adjacency matrix for message-passing, $\hat{\mathbf{A}} = \hat{\mathbf{A}}^{\text{asym}}$.

Lemma 3.1. *The expected propagation matrix under DropEdge (DE) is given as:*

$$\begin{aligned} \dot{\mathbf{P}}_{ii} &:= \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ii}] = \frac{1 - q^{d_i+1}}{(1-q)(d_i+1)} \\ \dot{\mathbf{P}}_{ij} &:= \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ij}] = \frac{1}{d_i} \left(1 - \frac{1 - q^{d_i+1}}{(1-q)(d_i+1)} \right) \end{aligned} \quad (3.3)$$

where $q \in [0, 1)$ is the dropping probability and $(j \rightarrow i) \in \mathcal{E}$.⁴

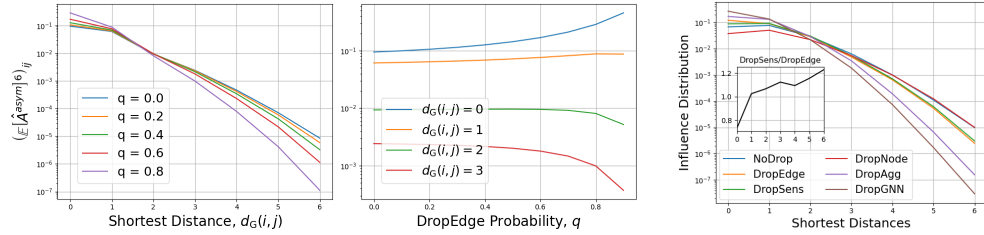
See Appendix B.1 for a proof, and a similar treatment of DropNode, DropAgg and DropGNN.

1-Layer Linear GCNs. $\forall q \in (0, 1)$ and nodes i, j such that $(j \rightarrow i) \in \mathcal{E}$, we have

$$\begin{aligned} \dot{\mathbf{P}}_{ii} &= \frac{1}{d_i + 1} \sum_{k=0}^{d_i} q^k > \frac{1}{d_i + 1} \\ \dot{\mathbf{P}}_{ij} &= \frac{1}{d_i} (1 - \dot{\mathbf{P}}_{ii}) < \frac{1}{d_i + 1} \end{aligned} \quad (3.4)$$

where the right-hand sides of the two inequalities are the corresponding entries in the propagation matrix of a NoDrop model. Equations 3.3, 3.4, B.16 and B.17 together imply the following result:

⁴We use $\dot{\mathbf{P}}$ to denote the expected propagation matrix under the asymmetric propagation rule. In Appendix D.2, we use $\ddot{\mathbf{P}}$ to denote the expected propagation matrix under the symmetric propagation rule.



(a) Entries of $\hat{P}^6$ decay at exponential rate w.r.t. distance between nodes, and polynomial rate w.r.t. to the DropEdge probability. (b) MC-approximation of influence distribution in ReLU-GCNs.

Figure 1: Empirical sensitivity analysis using the Cora dataset. The plot embedded in Figure 1b shows the ratio of influence under DropSens to that under DropEdge, at different distances.

Lemma 3.2. *In a 1-layer linear GCN with $\hat{A} = \hat{A}^{asym}$, using DropEdge, DropAgg or DropGNN*

1. *increases the sensitivity of a node's representations to its own input features, and*
2. *decreases the sensitivity to its neighbors' features.*

L-layer Linear GCNs. Unfortunately, we cannot draw similar conclusions in L-layer networks, for nodes at arbitrary distances. To see this, view $\hat{P}$ as the transition matrix of a non-uniform random walk. This walk has higher self-transition ($i = j$) probabilities than in a uniform augmented random walk ($T = \hat{A}^{asym}$, $q = 0$), but lower inter-node ($i \neq j$) transition probabilities. Note that $\hat{P}^L$ and T^L store the L-step transition probabilities in the corresponding walks. Then, since the paths connecting the nodes $i \in \mathcal{V}$ and $j \in \mathbb{B}^{(L-1)}(i)$ may involve self-loops, $(\hat{P}^L)_{ij}$ may be lower or higher than $(T^L)_{ij}$. Therefore, we cannot conclude how sensitivity between nodes separated by at most $L - 1$ hops changes. For nodes L-hops away, however, we can show that DropEdge always decreases the corresponding entry in $\hat{P}^L$, reducing the effective reachability of GCNs. Using Equations B.16 and B.17, we can show the same for DropAgg and DropGNN, respectively.

Theorem 3.1. *In an L-layer linear GCN with $\hat{A} = \hat{A}^{asym}$, using DropEdge, DropAgg or DropGNN decreases the sensitivity of a node $i \in \mathcal{V}$ to another node $j \in \mathbb{S}^{(L)}(i)$, thereby reducing its effective receptive field. Moreover, the sensitivity decreases with increasing dropping probability.*

See Appendix B.2 for a precise quantitative statement and the proof.

Nodes at Arbitrary Distances. Although no general statement could be made about the change in sensitivity between nodes up to $L - 1$ hops away, we can analyze such pairs empirically. We compute the L-hop transition matrix $\hat{P}^L$ – proportional to expected sensitivity in linear GCNs under DropEdge – for the Cora dataset, and average the entries after binning node pairs by the shortest distance between them. The results are shown in Figure 1a. In the left subfigure, we observe that the expected sensitivity decays at an *exponential rate* with increasing distance between the corresponding nodes. In the middle subfigure, we observe that DropEdge increases the expected sensitivity between nodes close to each other (0-hop and 1-hop neighbors) in the original topology, but reduces it between nodes farther off. Similar conclusions can be made with the symmetrically normalized propagation matrix (see Appendix D.2). Note that the over-squashing effects of DropAgg and DropGNN would, in theory, be even more severe, as suggested by Equations B.16 and B.17.

Nonlinear MPNNs. While linear networks are useful in simplifying the theoretical analysis, they are often not practical. In Appendix C.1, we treat the upper bounds on sensitivity established in previous works, and extend them to the DropEdge-variants. Even still, although theoretical bounds offer valuable guarantees, they can be arbitrarily loose in the absence of error quantification, making their practical relevance unclear. To reliably conclude the empirical behaviour of DropEdge- and Dropout-variants, we turn to Monte Carlo simulations with ReLU-GCNs; see Appendix D.1 for a description of the experiment setup. Figure 1b compares the influence of the source nodes [86] at different distances using a dropout probability of 0.5. We observe that while the effect of DropNode on the sensitivity profile – as compared to the baseline NoDrop – is relatively insignificant, models

using DropEdge, DropAgg and DropGNN have remarkably lower sensitivity to distant nodes, as predicted by our theory.

4 Sensitivity-Aware DropEdge

Lemma 3.1 tells us that DropEdge decreases the weight of cross-edges, $(j \rightarrow i)$, in the expected propagation matrix, i.e. the *strength of message passing* over these edges decreases. The fraction of *information preserved* over a cross-edge is dependent only on the dropping probability and the target node's in-degree, d_i . We can directly control this quantity using a per-edge dropping probability, q_i , dependent only on the receiving node's in-degree:

$$c := \frac{\mathbb{E}_{\text{DE}}[\hat{A}_{ij}]}{\mathbb{E}_{\text{ND}}[\hat{A}_{ij}]} = \frac{d_i + 1}{d_i} \left(1 - \frac{1 - q_i^{d_i+1}}{(1 - q_i)(d_i + 1)} \right) \implies 1 - c = \frac{q_i - q_i^{d_i+1}}{d_i(1 - q_i)} \quad (4.1)$$

where c is the fraction of information preserved, e.g. 95%, and ND refers to a NoDrop model. We can solve for q_i and mask the incoming edges to node i accordingly; we name this algorithm *DropSens*. In [Appendix E.3](#), we present a Python implementation of the algorithm, as well as a computationally efficient approximation to [Equation 4.1](#). In the plot embedded in [Figure 1b](#), we can observe that DropSens increases the influence of distant nodes, as compared to DropEdge.

Note that higher values of c encourage lower q_i , while lower values permit q_i to take higher values; see [Figure 7a](#). Since this can result in abnormally high dropping probabilities, we clip the value of q_i in our experiments using another hyperparameter, $q_{\max}$; see details in [Appendix E.2](#).

Challenges in Extending to GIN and GAT. DropSens has been specifically derived for GCN-style architectures, where the edge weights used in message aggregation are simple functions of node degree. Extending it to other architectures raises nontrivial challenges:

1. GIN uses constant edge weights, so the sensitivity of a node to its neighbors simplifies to $(1 - q) \|\mathbf{W}\|$ for dropping probability q . Enforcing a fixed information preservation ratio c gives $q = 1 - c$, which is equivalent to DropEdge. This limitation arises because GIN's neighbour-weighting scheme is independent of the node degrees, making a principled variant of DropSens unnecessary.
2. GAT, in contrast, uses feature-dependent attention weights that vary at every layer and iteration. This means a DropSens-style approach would require recomputing edge masks in each iteration and each layer, compromising the simplicity we aimed for. Moreover, the presence of softmax attention makes the closed-form of sensitivity complicated.

5 Experiments

Our theoretical analysis indicates that random dropping may degrade the performance of GNNs in tasks that depend on capturing LRIs. In this section, we test this hypothesis by evaluating DropEdge- and Dropout-variants on both synthetic and real-world datasets. A complete description of the datasets is provided in [Appendix E.1](#), and the experimental details are in [Appendix E.2](#).

5.1 Synthetic Datasets

SyntheticZINC [33] is a synthetic variant of the ZINC dataset [42], designed to study the effect of information mixing in graph learning. Node features are sparsely assigned, and the target requires non-linear mixing of two selected nodes' features, chosen based on their commute time [10]. We vary the mixing level and evaluate an 11-layer GCN, ensuring sufficient message passing. For better readability, we only test DropEdge, Dropout and DropMessage – the three more popular methods used for training deep GNNs.

The results are presented in [Figure 2](#), where we can observe that the mean absolute error (MAE) increases with the commute time percentile used to select the node pairs, as was hypothesized and evidenced in [33]. Additionally, we observe that both train and test performance decline when using dropout with a probability as low as 0.2, and even more so with a higher probability of 0.5. These results provide strong evidence for the detrimental effects of dropout methods in modelling long-range interactions, supporting our theoretical analysis.

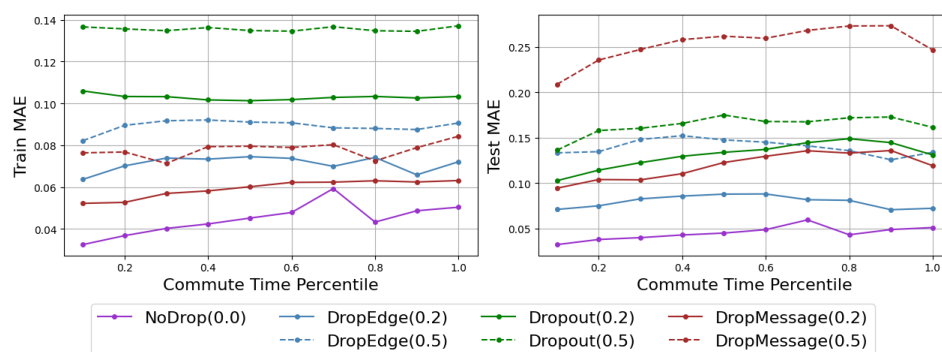


Figure 2: Train and test MAE of 11-layer GCNs on the SyntheticZINC dataset, averaged over 10 initializations.

5.2 Real-world Datasets

To test the dropping methods on real-world datasets, we use the GCN, GIN [87] and GAT [83] architectures – GCN and GIN satisfy the model assumptions made in all the theoretical results presented in Section 3, while GAT does not satisfy any of them, since the attention scores are computed as a function of *all* the node representations. Therefore, GCN, GIN and GAT together provide a broad representation of different MPNN architectures. We present the results for GCN and GIN in the main text, since these models were used as baselines in a majority of works on alleviating over-squashing [4, 5, 8, 35, 44, 69, 81]; the results for GAT are reported in Table 7.

For each dataset–model–dropout combination, we perform 20 independent runs to find the best performing dropout configuration; results are reported in Table 10. We then perform a t-test to assess whether dropout improves performance, using 50 samples from the NoDrop model ($q = 0$) and 50 samples from the best performing dropout configuration.⁵ In this section, we report the p-values of the tests, and in Table 8, we report the effect sizes as Hedges’ g statistic [39].

Node-classification. Although determining whether a task requires modelling LRIs can be challenging, understanding the structure of the datasets can provide important insight. For example, homophilic datasets have local consistency in node labels, i.e. nodes closely connected to each other have similar labels. On the other hand, in heterophilic datasets, nearby nodes often have dissimilar labels. Since DropEdge-variants increase the sensitivity of a node’s representations to its immediate neighbors, and reduce its sensitivity to distant nodes, we expect it to improve performance on homophilic datasets but harm performance on heterophilic ones; such a setup was also used in [81]. In this work, we use Cora [58], CiteSeer [31] and PubMed [63] as representatives of homophilic datasets [53, 100], and Squirrel, Chameleon and TwitchDE [71] to represent heterophilic datasets [53]. The networks’ statistics are presented in Table 4, where we can note the remarkably lower homophily measures of heterophilic datasets.

The results are presented in Table 1a. It is clear to see that dropout *can significantly improve* test performance on homophilic datasets – with $31/54 \approx 57\%$ cases performing better than the corresponding NoDrop baseline – indicating that these methods are indeed beneficial in tackling short-range tasks. On the other hand, with the heterophilic datasets, the improvement is *insignificant*. Rather, in most ($46/54 \approx 85\%$) cases, the best dropout configuration performs worse than the NoDrop baseline. This suggests that the dropping methods harm generalization in long-range tasks by forcing models to overfit to short-range signals (see Appendix F.3 for supporting evidence).

Graph-classification. Several graph classification datasets have also been identified as long-range tasks, like the molecular networks datasets Mutag [17], Proteins [19] and Enzymes [9], and the social networks datasets Reddit, IMDb and Collab [89]. These datasets have also been used for evaluation in previous works on over-squashing, including [8, 44].

⁵The t-test assumes that both samples are drawn from normal distributions – all Shapiro-Wilk tests for non-normality of samples [76] failed at 90% confidence.

Table 1: Difference in mean test accuracy (%) between the best performing configuration of each dropout method and the baseline NoDrop model. Cell colors represent p-values from a t-test evaluating whether dropout improves performance: green indicates significance at 90% confidence, while red denotes insignificant results.

(a) Node classification tasks.

GNN	Dropout	Homophilic Networks			Heterophilic Networks		
		Cora	CiteSeer	PubMed	Chameleon	Squirrel	TwitchDE
GCN	DropEdge	+0.215	+0.515	−0.100	−1.711	−0.334	−0.041
	DropNode	+0.213	−0.302	−0.082	−4.987	−1.307	−1.180
	DropAgg	+0.092	+0.552	−0.429	−15.180	−13.738	−5.306
	DropGNN	+0.372	+0.727	−0.194	−3.841	−1.451	−0.950
	Dropout	+0.418	−0.262	+1.111	−3.672	−1.298	−0.516
	DropMessage	+0.140	+0.092	+0.820	−1.501	−0.549	+0.270
GIN	DropEdge	−5.370	−11.911	−3.977	−7.814	−1.503	+0.724
	DropNode	+3.983	+4.421	+1.260	−2.655	−0.871	+0.426
	DropAgg	−2.465	−11.098	−2.038	−7.233	−4.973	−0.928
	DropGNN	−0.929	−10.150	−2.280	−9.544	−5.118	−0.057
	Dropout	+4.771	+3.128	+2.311	−0.865	+1.062	+1.254
	DropMessage	+5.661	+2.841	+2.003	+4.243	+2.439	+1.919

(b) Graph classification tasks.

GNN	Dropout	Molecular Networks			Social Networks		
		Mutag	Proteins	Enzymes	Reddit	IMDb	Collab
GCN	DropEdge	−1.200	+0.482	−3.032	−6.280	+0.680	−3.353
	DropNode	+0.000	+0.875	−3.083	−10.930	−0.320	−4.511
	DropAgg	+1.100	−0.286	−5.424	−17.450	+0.540	−20.940
	DropGNN	+2.600	+0.482	−2.900	−12.940	−1.180	+0.880
	Dropout	−2.700	−0.089	−4.485	−0.920	+0.580	−3.299
	DropMessage	−0.200	+0.393	−3.584	−3.770	+3.420	+2.607
GIN	DropEdge	−2.900	−1.143	−8.954	−0.430	−0.240	−0.558
	DropNode	−8.500	+0.536	−7.568	−1.500	+0.380	+0.125
	DropAgg	+0.200	+0.143	−0.640	−1.560	−0.800	−0.988
	DropGNN	−3.700	−0.661	−11.745	−4.150	−2.740	−3.456
	Dropout	+0.600	+0.661	−2.042	+0.390	+0.080	+0.383
	DropMessage	+1.100	+1.536	−2.406	−0.130	+1.340	+1.506

The results are shown in Table 1b, where we observe that dropout methods are usually detrimental to model performance, often recording a non-positive effect (62/102 $\approx$ 61% cases).⁶ Notably, the p-values are lower as compared to those recorded for heterophilic datasets in Table 1a, i.e. higher evidence for efficacy of dropping methods. We conjecture that over-squashing may have limited impact on model performance in graph-level tasks since the aggregation module eventually mixes information from distant nodes for computing graph-level representations.

5.3 Evaluating DropSens

We start by comparing DropSens with DropEdge on real-world datasets from Section 5.2, illustrating how algorithms can be readily adapted for better suitability at modelling LRIs. In Figure 3, we present the relative change in error rate (1 − Acc) of DropSens w.r.t. DropEdge, with GCN as the base model. It is clear to observe a uniform improvement in the performance on long-range tasks, suggesting that addressing over-squashing using DropSens can enhance the effectiveness of GCNs.

Finally, we benchmark DropSens against state-of-the-art graph-rewiring techniques designed specifically to tackle over-squashing (see Appendix A.3 for their descriptions). We train a GCN on node classification tasks, following the setup in [81], and both GCN and GIN on graph classification tasks, following [8, 44]. The results with GCN are reported in Table 2, where we find that DropSens

⁶We exclude Collab×GAT due to unsuccessful runs.

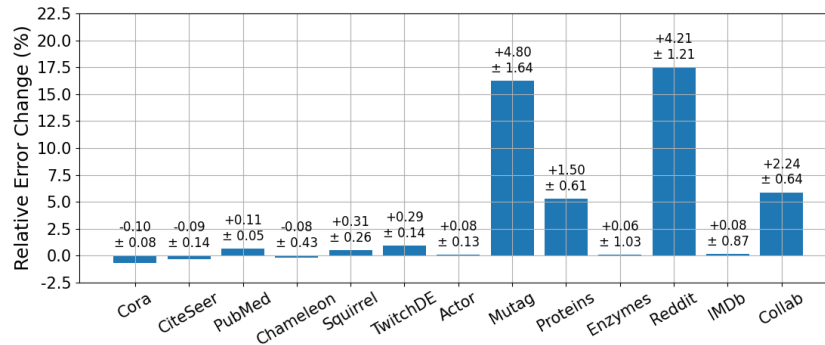


Figure 3: Relative change in test-time performance of a GCN using DropSens, compared to the baseline DropEdge, on real-world datasets from Section 5.2. Text over each bar indicates the difference between mean performance of DropSens and DropEdge, and corresponding standard errors.

Table 2: Performance of GCN with graph rewiring methods. **First**, **second**, and **third** best results are coloured.

(a) Node-classification tasks – results for other methods taken from [81, Table 2].

Rewiring	Cora	CiteSeer	PubMed	Chameleon	Squirrel	Actor
None	81.89	72.31	78.16	41.33	30.32	23.84
Undirected	-	-	-	42.02	35.53	21.45
+FA	81.65	70.47	79.48	42.67	36.86	24.14
DIGL (PPR)	83.21	73.29	78.84	42.02	33.22	24.77
DIGL + Undirected	-	-	-	42.68	32.48	25.45
SDRF	82.76	72.58	79.10	42.73	37.05	28.42
SDRF + Undirected	-	-	-	44.46	37.67	28.35
DropSens	85.17	73.81	83.76	53.45	39.72	22.89

(b) Graph-classification tasks – results for other methods from [8, Table 1].

Rewiring	Mutag	Proteins	Enzymes	Reddit	IMDb	Collab
None	72.15	70.98	27.67	68.26	49.77	33.78
Last FA	70.05	71.02	26.47	68.49	48.98	33.32
Every FA	70.45	60.04	18.33	48.49	48.17	51.80
DIGL	79.70	70.76	35.72	76.04	64.39	54.50
SDRF	71.05	70.92	28.37	68.62	49.40	33.45
FoSR	80.00	73.42	25.07	70.33	49.66	33.84
GTR	79.10	72.59	27.52	68.99	49.92	33.05
DropSens	75.30	73.27	28.42	80.16	49.34	64.16

outperforms other methods in node classification tasks, and performs competitively in graph classification tasks. In addition to superior performance, another advantage of DropSens over the other methods is that it significantly reduces the number of messages being propagated, thereby tackling the problem of over-smoothing and increasing training speed.

The results with GIN are presented in Appendix G.4, where we observe that DropSens does not perform competitively – unsurprising, since DropSens was specifically designed to work with GCN’s message-passing scheme.

6 Conclusion

There exists an important gap in our understanding of several algorithms designed for training deep GNNs – while their positive effects on model performance have been well-studied, making them popular choices for training deep GNNs, their evaluation has been limited to short-range tasks.

This is rooted in a key assumption: that if a deep GNN is trainable, it must also be capable of modelling LRIs. As a result, potential adverse effects of these algorithms on capturing LRIs have been overlooked. Our results challenge this assumption – we theoretically and empirically show that DropEdge- and Dropout-variants exacerbate the over-squashing problem in deep GNNs, and degrade performance on long-range tasks. This highlights the need for a more comprehensive evaluation of common training practices for deep GNNs, with special emphasis on their capacity to capture LRIs. This is crucial for building confidence in their use beyond controlled benchmarks.

Limitations. While our theoretical analysis successfully predicts how DropEdge-variants affect test performance on short-range and long-range tasks, it is based on several simplifying assumptions on the message-passing scheme. These assumptions, although standard in the literature, limit the generalizability of our conclusions to other architectures, including ResGCNs [51], GATs [83], and Graph Transformers [85]. Additionally, an important limitation of DropSens is that it requires an architecture-specific alteration to the edge-dropping strategy, which is not practical in general. As also mentioned in Section 4, we did not intend to introduce DropSens as a benchmark, but rather to demonstrate how methods designed for alleviating over-smoothing can be readily adapted to simultaneously control over-squashing.

Future Directions. Currently, real-world datasets are classified as short- or long-range tasks based on extensive model training [4] or weak proxy measures like node homophily [81]. Developing a reliable measure of information mixing in the ground-truth data could greatly benefit the research community. Such a measure would enable more precise identification of short-, intermediate- and long-range tasks, improving evaluation and benchmarking. Another interesting direction is to investigate the significance of over-squashing in graph-level tasks, where the aggregation module of MPNNs enables some mixing of information from distant nodes. To the best of our knowledge, [33] is the only work that directly addresses this question, offering strong theoretical insights. However, empirical validation of these effects remains limited.

Acknowledgement

Keyue Jiang was supported by the UKRI Engineering and Physical Sciences Research Council (EPSRC) [grant number EP/R513143/1]. The authors would like to thank the anonymous reviewers for their helpful comments and insights.

References

- [1] Zeyuan Allen-Zhu, Aditya Bhaskara, Silvio Lattanzi, Vahab Mirrokni, and Lorenzo Orecchia. Expanders via local edge flips. In *Proceedings of the Twenty-Seventh Annual ACM-SIAM Symposium on Discrete Algorithms*, SODA '16, pp. 259–269, USA, 2016. Society for Industrial and Applied Mathematics.
- [2] N. Alon and V. D. Milman. λ_1 , isoperimetric inequalities for graphs, and superconcentrators. *Journal of Combinatorial Theory, Series B*, 38(1):73–88, February 1985.
- [3] Noga Alon. Eigenvalues and expanders. *Combinatorica*, 6(2):83–96, June 1986.
- [4] Uri Alon and Eran Yahav. On the bottleneck of graph neural networks and its practical implications. In *International Conference on Learning Representations*, 2021.
- [5] Adrián Arnaiz-Rodríguez, Ahmed Begga, Francisco Escolano, and Nuria M Oliver. Diffwire: Inductive graph rewiring via the lovász bound. In Bastian Rieck and Razvan Pascanu (eds.), *Proceedings of the First Learning on Graphs Conference*, volume 198 of *Proceedings of Machine Learning Research*, pp. 15:1–15:27. PMLR, 12 2022.
- [6] Pradeep Kr. Banerjee, Kedar Karhadkar, Yu Guang Wang, Uri Alon, and Guido Montúfar. Oversquashing in gnns through the lens of information contraction and graph expansion. In *2022 58th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, pp. 1–8. IEEE Press, 2022.
- [7] Pablo Barceló, Egor V. Kostylev, Mikael Monet, Jorge Pérez, Juan Reutter, and Juan Pablo Silva. The logical expressiveness of graph neural networks. In *International Conference on Learning Representations*, 2020.

- [8] Mitchell Black, Zhengchao Wan, Amir Nayyeri, and Yusu Wang. Understanding oversquashing in GNNs through the lens of effective resistance. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett (eds.), *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pp. 2528–2547. PMLR, 07 2023.
- [9] Karsten M. Borgwardt, Cheng Soon Ong, Stefan Schönauer, S. V. N. Vishwanathan, Alex J. Smola, and Hans-Peter Kriegel. Protein function prediction via graph kernels. *Bioinformatics*, 21(suppl_1):i47–i56, 06 2005.
- [10] Ashok K. Chandra, Prabhakar Raghavan, Walter L. Ruzzo, Roman Smolensky, and Prasoon Tiwari. The electrical resistance of a graph captures its commute and cover times. *computational complexity*, 6:312–340, 1989.
- [11] Jianfei Chen, Jun Zhu, and Le Song. Stochastic training of graph convolutional networks with variance reduction. In *International Conference on Machine Learning*, pp. 941–949, 2018.
- [12] Jie Chen, Tengfei Ma, and Cao Xiao. FastGCN: Fast learning with graph convolutional networks via importance sampling. In *International Conference on Learning Representations*, 2018.
- [13] Ming Chen, Zhewei Wei, Zengfeng Huang, Bolin Ding, and Yaliang Li. Simple and deep graph convolutional networks. In Hal Daumé III and Aarti Singh (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 1725–1735. PMLR, 07 2020.
- [14] Jacob Cohen. *Statistical Power Analysis for the Behavioral Sciences* statistical power analysis for the behavioral sciences. Lawrence Erlbaum Associates, Hillsdale, NJ, 2 edition, 1988.
- [15] Colin Cooper, Martin Dyer, Catherine Greenhill, and Andrew Handley. The flip markov chain for connected regular graphs. *Discrete Applied Mathematics*, 254:56–79, 2019.
- [16] Andreea Deac, Marc Lackenby, and Petar Veličković. Expander graph propagation. In *NeurIPS 2022 Workshop: New Frontiers in Graph Learning*, 2022.
- [17] Asim Kumar Debnath, Rosa L. Lopez de Compadre, Gargi Debnath, Alan J. Shusterman, and Corwin Hansch. Structure-activity relationship of mutagenic aromatic and heteroaromatic nitro compounds. correlation with molecular orbital energies and hydrophobicity. *Journal of Medicinal Chemistry*, 34(2):786–797, Feb 1991.
- [18] Francesco Di Giovanni, Lorenzo Giusti, Federico Barbero, Giulia Luise, Pietro Lio, and Michael M Bronstein. On over-squashing in message passing neural networks: The impact of width, depth, and topology. In *International Conference on Machine Learning*, pp. 7865–7885. PMLR, 2023.
- [19] Paul D. Dobson and Andrew J. Doig. Distinguishing enzyme structures from non-enzymes without alignments. *Journal of Molecular Biology*, 330(4):771–783, 2003.
- [20] Vijay Prakash Dwivedi, Ladislav Rampásek, Mikhail Galkin, Ali Parviz, Guy Wolf, Anh Tuan Luu, and Dominique Beaini. Long range graph benchmark. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022.
- [21] Taoran Fang, Zhiqing Xiao, Chunping Wang, Jiarong Xu, Xuan Yang, and Yang Yang. Dropmessage: Unifying random dropping for graph neural networks. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(4):4267–4275, Jun. 2023.
- [22] Tomas Feder, Adam Guetz, Milena Mihail, and Amin Saberi. A local switch markov chain on given degree graphs with application in connectivity of peer-to-peer networks. In *2006 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS'06)*, pp. 69–76, 2006.
- [23] Wenzheng Feng, Jie Zhang, Yuxiao Dong, Yu Han, Huanbo Luan, Qian Xu, Qiang Yang, Evgeny Kharlamov, and Jie Tang. Graph random neural networks for semi-supervised learning on graphs. In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 22092–22103. Curran Associates, Inc., 2020.

- [24] Matthias Fey and Jan E. Lenssen. Fast graph representation learning with PyTorch Geometric. In *ICLR Workshop on Representation Learning on Graphs and Manifolds*, 2019.
- [25] Rickard Br  el Gabrielsson, Mikhail Yurochkin, and Justin Solomon. Rewiring with positional encodings for graph neural networks. *Transactions on Machine Learning Research*, 2023.
- [26] Yarin Gal and Zoubin Ghahramani. Bayesian convolutional neural networks with bernoulli approximate variational inference, 2016.
- [27] Yarin Gal and Zoubin Ghahramani. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In Maria Florina Balcan and Kilian Q. Weinberger (eds.), *Proceedings of The 33rd International Conference on Machine Learning*, volume 48 of *Proceedings of Machine Learning Research*, pp. 1050–1059, New York, New York, USA, 06 2016. PMLR.
- [28] Yarin Gal and Zoubin Ghahramani. A theoretically grounded application of dropout in recurrent neural networks. In D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 29. Curran Associates, Inc., 2016.
- [29] Hongyang Gao, Zhengyang Wang, and Shuiwang Ji. Large-scale learnable graph convolutional networks. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 1416–1424. ACM, 2018.
- [30] George Giakkoupis. Expanders via local edge flips in quasilinear time. In *Proceedings of the 54th Annual ACM SIGACT Symposium on Theory of Computing*, STOC 2022, pp. 64–76, New York, NY, USA, 2022. Association for Computing Machinery.
- [31] C. Lee Giles, Kurt D. Bollacker, and Steve Lawrence. Citeseer: an automatic citation indexing system. In *Proceedings of the Third ACM Conference on Digital Libraries*, DL ’98, pp. 89–98, New York, NY, USA, 1998. Association for Computing Machinery.
- [32] Justin Gilmer, Samuel S. Schoenholz, Patrick F. Riley, Oriol Vinyals, and George E. Dahl. Neural message passing for quantum chemistry. In Doina Precup and Yee Whye Teh (eds.), *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pp. 1263–1272. PMLR, 08 2017.
- [33] Francesco Di Giovanni, T. Konstantin Rusch, Michael Bronstein, Andreea Deac, Marc Lackenby, Siddhartha Mishra, and Petar Veli  kovi  . How does over-squashing affect the power of GNNs? *Transactions on Machine Learning Research*, 2024.
- [34] Jhony H. Giraldo, Konstantinos Skianis, Thierry Bouwmans, and Fragkiskos D. Malliaros. On the trade-off between over-smoothing and over-squashing in deep graph neural networks. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management*, CIKM ’23, pp. 566–576, New York, NY, USA, 2023. Association for Computing Machinery.
- [35] Benjamin Gutteridge, Xiaowen Dong, Michael M Bronstein, and Francesco Di Giovanni. DRew: Dynamically rewired message passing with delay. In *International Conference on Machine Learning*, pp. 12252–12267. PMLR, 2023.
- [36] Will Hamilton, Zhitao Ying, and Jure Leskovec. Inductive representation learning on large graphs. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [37] Arman Hasanzadeh, Ehsan Hajiramezanali, Shahin Boluki, Mingyuan Zhou, Nick Duffield, Krishna Narayanan, and Xiaoning Qian. Bayesian graph neural networks with adaptive connection sampling, 2020.
- [38] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 770–778, 2016.

- [39] Larry V. Hedges. Distribution theory for glass's estimator of effect size and related estimators. *Journal of Educational Statistics*, 6(2):107–128, 2025/03/15/ 1981. Full publication date: Summer, 1981.
- [40] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. Open graph benchmark: Datasets for machine learning on graphs. *arXiv preprint arXiv:2005.00687*, 2020.
- [41] Lei Huang, Jie Qin, Yi Zhou, Fan Zhu, Li Liu, and Ling Shao. Normalization techniques in training dnns: Methodology, analysis and application, 2020.
- [42] John J Irwin, Teague Sterling, Michael M Mysinger, Erin S Bolstad, and Ryan G Coleman. ZINC: a free tool to discover chemistry for biology. *J Chem Inf Model*, 52(7):1757–1768, June 2012.
- [43] Bo Jiang, Yong Chen, Beibei Wang, Haiyun Xu, and Bin Luo. Dropagg: Robust graph neural networks via drop aggregation. *Neural Networks*, 163:65–74, 2023.
- [44] Kedar Karhadkar, Pradeep Kr. Banerjee, and Guido Montufar. FoSR: First-order spectral rewiring for addressing oversquashing in GNNs. In *The Eleventh International Conference on Learning Representations*, 2023.
- [45] Kenji Kawaguchi. Deep learning without poor local minima. In D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 29. Curran Associates, Inc., 2016.
- [46] Diederik Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations (ICLR)*, San Diego, CA, USA, 2015.
- [47] Thomas N. Kipf and Max Welling. Semi-supervised classification with graph convolutional networks. In *International Conference on Learning Representations (ICLR)*, 2017.
- [48] Yann LeCun, Bernhard Boser, John Denker, Donnie Henderson, R. Howard, Wayne Hubbard, and Lawrence Jackel. Handwritten digit recognition with a back-propagation network. In D. Touretzky (ed.), *Advances in Neural Information Processing Systems*, volume 2. Morgan-Kaufmann, 1989.
- [49] Guohao Li, Matthias Müller, Ali Thabet, and Bernard Ghanem. Deepgcns: Can gcns go as deep as cnns? In *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 9266–9275, 2019.
- [50] Qimai Li, Zhichao Han, and Xiao-ming Wu. Deeper insights into graph convolutional networks for semi-supervised learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1), 04 2018.
- [51] Yujia Li, Daniel Tarlow, Marc Brockschmidt, and Richard S. Zemel. Gated graph sequence neural networks. In Yoshua Bengio and Yann LeCun (eds.), *4th International Conference on Learning Representations, ICLR 2016, San Juan, Puerto Rico, May 2-4, 2016, Conference Track Proceedings*, 2016.
- [52] Huidong Liang, Haitz Sáez de Ocáriz Borde, Baskaran Sripathamathan, Michael Bronstein, and Xiaowen Dong. Towards quantifying long-range interactions in graph machine learning: a large graph dataset and a measurement, 2025.
- [53] Derek Lim, Xiuyu Li, Felix Hohne, and Ser-Nam Lim. New benchmarks for learning on non-homophilous graphs. *arXiv preprint arXiv:2104.01404*, 2021.
- [54] Meng Liu, Hongyang Gao, and Shuiwang Ji. Towards deeper graph neural networks. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. ACM, 2020.
- [55] Yang Liu, Chuan Zhou, Shirui Pan, Jia Wu, Zhao Li, Hongyang Chen, and Peng Zhang. Curvdop: A ricci curvature based approach to prevent graph neural networks from over-smoothing and over-squashing. In *Proceedings of the ACM Web Conference 2023*, WWW '23, pp. 221–230, New York, NY, USA, 2023. Association for Computing Machinery.

- [56] L. Lovász. Random walks on graphs: A survey. *Combinatorics, Paul Erdos is Eighty*, 2(1): 1–46, 1993.
- [57] Peter Mahlmann and Christian Schindelhauer. Peer-to-peer networks based on random transformations of connected regular undirected graphs. In *Proceedings of the Seventeenth Annual ACM Symposium on Parallelism in Algorithms and Architectures*, SPAA '05, pp. 155–164, New York, NY, USA, 2005. Association for Computing Machinery.
- [58] Andrew Kachites McCallum, Kamal Nigam, Jason Rennie, and Kristie Seymore. Automating the construction of internet portals with machine learning. *Information Retrieval*, 3(2):127–163, 07 2000.
- [59] Péter Mernyei and Cătălina Cangea. Wiki-cs: A wikipedia-based benchmark for graph neural networks. *arXiv preprint arXiv:2007.02901*, 2020.
- [60] Aaron Meurer, Christopher P. Smith, Mateusz Paprocki, Ondřej Čertík, Sergey B. Kirpichev, Matthew Rocklin, AMiT Kumar, Sergiu Ivanov, Jason K. Moore, Sartaj Singh, Thilina Rathnayake, Sean Vig, Brian E. Granger, Richard P. Muller, Francesco Bonazzi, Harsh Gupta, Shivam Vats, Fredrik Johansson, Fabian Pedregosa, Matthew J. Curry, Andy R. Terrel, Štěpán Roučka, Ashutosh Saboo, Isuru Fernando, Sumith Kulal, Robert Cimrman, and Anthony Scopatz. Sympy: symbolic computing in python. *PeerJ Computer Science*, 3:e103, January 2017.
- [61] Federico Monti, Michael Bronstein, and Xavier Bresson. Geometric matrix completion with recurrent multi-graph neural networks. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [62] Christopher Morris, Nils M. Kriege, Franka Bause, Kristian Kersting, Petra Mutzel, and Marion Neumann. Tudataset: A collection of benchmark datasets for learning with graphs. In *ICML 2020 Workshop on Graph Representation Learning and Beyond (GRL+ 2020)*, 2020.
- [63] Galileo Mark Namata, Ben London, Lise Getoor, and Bert Huang. Query-driven active surveying for collective classification. In *International Workshop on Mining and Learning with Graphs*, Edinburgh, Scotland, 2012. MLG.
- [64] Khang Nguyen, Hieu Nong, Vinh Nguyen, Nhat Ho, Stanley Osher, and Tan Nguyen. Revisiting over-smoothing and over-squashing using ollivier-ricci curvature. In *Proceedings of the 40th International Conference on Machine Learning, ICML'23*. JMLR.org, 2023.
- [65] Kenta Oono and Taiji Suzuki. Graph neural networks exponentially lose expressive power for node classification. In *International Conference on Learning Representations*, 2020.
- [66] Pál András Papp, Karolis Martinkus, Lukas Faber, and Roger Wattenhofer. Dropgnn: Random dropouts increase the expressiveness of graph neural networks. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan (eds.), *Advances in Neural Information Processing Systems*, volume 34, pp. 21997–22009. Curran Associates, Inc., 2021.
- [67] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Köpf, Edward Yang, Zach 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, 2019.
- [68] Hongbin Pei, Bingzhe Wei, Kevin Chen-Chuan Chang, Yu Lei, and Bo Yang. Geom-gcn: Geometric graph convolutional networks. In *International Conference on Learning Representations*, 2020.
- [69] Chendi Qian, Andrei Manolache, Kareem Ahmed, Zhe Zeng, Guy Van den Broeck, Mathias Niepert, and Christopher Morris. Probabilistically rewired message-passing neural networks. In *The Twelfth International Conference on Learning Representations*, 2024.

- [70] Yu Rong, Wenbing Huang, Tingyang Xu, and Junzhou Huang. Dropedge: Towards deep graph convolutional networks on node classification. In *International Conference on Learning Representations*, 2020.
- [71] Benedek Rozemberczki, Carl Allen, and Rik Sarkar. Multi-Scale Attributed Node Embedding. *Journal of Complex Networks*, 9(2), 2021.
- [72] T. Konstantin Rusch, Michael M. Bronstein, and Siddhartha Mishra. A survey on oversmoothing in graph neural networks, 2023.
- [73] Franco Scarselli, Marco Gori, Ah Chung Tsoi, Markus Hagenbuchner, and Gabriele Monfardini. The graph neural network model. *IEEE Transactions on Neural Networks*, 20(1):61–80, 2009.
- [74] Michael Schlichtkrull, Thomas N. Kipf, Peter Bloem, Rianne van den Berg, Ivan Titov, and Max Welling. Modeling relational data with graph convolutional networks, 2017.
- [75] Ida Schomburg, Antje Chang, Christian Ebeling, Marion Gremse, Christian Heldt, Gregor Huhn, and Dietmar Schomburg. BRENDA, the enzyme database: updates and major new developments. *Nucleic Acids Res*, 32(Database issue):D431–3, January 2004.
- [76] S. S. Shapiro and M. B. Wilk. An analysis of variance test for normality (complete samples). *Biometrika*, 52(3/4):591–611, 2025/03/12/ 1965. Full publication date: Dec., 1965.
- [77] Oleksandr Shchur, Maximilian Mumme, Aleksandar Bojchevski, and Stephan Günnemann. Pitfalls of graph neural network evaluation. *Relational Representation Learning Workshop, NeurIPS 2018*, 2018.
- [78] Alistair Sinclair and Mark Jerrum. Approximate counting, uniform generation and rapidly mixing markov chains. *Information and Computation*, 82(1):93–133, 1989.
- [79] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. Dropout: A simple way to prevent neural networks from overfitting. *Journal of Machine Learning Research*, 15(56):1929–1958, 2014.
- [80] Christian Szegedy, Wei Liu, Yangqing Jia, Pierre Sermanet, Scott Reed, Dragomir Anguelov, Dumitru Erhan, Vincent Vanhoucke, and Andrew Rabinovich. Going deeper with convolutions. In *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 1–9, 2015.
- [81] Jake Topping, Francesco Di Giovanni, Benjamin Paul Chamberlain, Xiaowen Dong, and Michael M. Bronstein. Understanding over-squashing and bottlenecks on graphs via curvature. In *International Conference on Learning Representations*, 2022.
- [82] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [83] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph attention networks. In *International Conference on Learning Representations*, 2018.
- [84] Nikil Wale and George Karypis. Comparison of descriptor spaces for chemical compound retrieval and classification. In *Sixth International Conference on Data Mining (ICDM'06)*, pp. 678–689, 2006.
- [85] Zhanghao Wu, Paras Jain, Matthew Wright, Azalia Mirhoseini, Joseph E Gonzalez, and Ion Stoica. Representing long-range context for graph neural networks with global attention. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.

- [86] Keyulu Xu, Chengtao Li, Yonglong Tian, Tomohiro Sonobe, Ken-ichi Kawarabayashi, and Stefanie Jegelka. Representation learning on graphs with jumping knowledge networks. In Jennifer Dy and Andreas Krause (eds.), *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pp. 5453–5462. PMLR, 07 2018.
- [87] Keyulu Xu, Weihua Hu, Jure Leskovec, and Stefanie Jegelka. How powerful are graph neural networks?, 2019.
- [88] Han Xuanyuan, Tianxiang Zhao, and Dongsheng Luo. Shedding light on random dropping and oversmoothing. In *NeurIPS 2023 Workshop: New Frontiers in Graph Learning*, 2023.
- [89] Pinar Yanardag and S.V.N. Vishwanathan. Deep graph kernels. In *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '15, pp. 1365–1374, New York, NY, USA, 2015. Association for Computing Machinery.
- [90] Rex Ying, Ruining He, Kaifeng Chen, Pong Eksombatchai, William L. Hamilton, and Jure Leskovec. Graph convolutional neural networks for web-scale recommender systems. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, KDD '18, pp. 974–983, New York, NY, USA, 2018. Association for Computing Machinery.
- [91] Yuning You, Tianlong Chen, Yongduo Sui, Ting Chen, Zhangyang Wang, and Yang Shen. Graph contrastive learning with augmentations. In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 5812–5823. Curran Associates, Inc., 2020.
- [92] Yuning You, Tianlong Chen, Zhangyang Wang, and Yang Shen. When does self-supervision help graph convolutional networks? In Hal Daumé III and Aarti Singh (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 10871–10880. PMLR, 07 2020.
- [93] Yuning You, Tianlong Chen, Zhangyang Wang, and Yang Shen. L2-gcn: Layer-wise and learned efficient training of graph convolutional networks. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 2124–2132, 2020.
- [94] Yuning You, Tianlong Chen, Zhangyang Wang, and Yang Shen. Bringing your own view: Graph contrastive learning without prefabricated data augmentations. In *Proceedings of the Fifteenth ACM International Conference on Web Search and Data Mining*, WSDM '22, pp. 1300–1309, New York, NY, USA, 2022. Association for Computing Machinery.
- [95] Lingxiao Zhao and Leman Akoglu. Pairnorm: Tackling oversmoothing in gnns. In *International Conference on Learning Representations*, 2020.
- [96] Wenqing Zheng, Edward W Huang, Nikhil Rao, Sumeet Katariya, Zhangyang Wang, and Karthik Subbian. Cold brew: Distilling graph node representations with incomplete or missing neighborhoods. In *International Conference on Learning Representations*, 2022.
- [97] Kaixiong Zhou, Xiao Huang, Yuening Li, Daochen Zha, Rui Chen, and Xia Hu. Towards deeper graph neural networks with differentiable group normalization. In *Advances in neural information processing systems*, 2020.
- [98] Kaixiong Zhou, Xiao Huang, Daochen Zha, Rui Chen, Li Li, Soo-Hyun Choi, and Xia Hu. Dirichlet energy constrained learning for deep graph neural networks. *Advances in neural information processing systems*, 2021.
- [99] Kuangqi Zhou, Yanfei Dong, Kaixin Wang, Wee Sun Lee, Bryan Hooi, Huan Xu, and Jia-ashi Feng. Understanding and resolving performance degradation in deep graph convolutional networks. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, pp. 2728–2737, 2021.

- [100] Jiong Zhu, Yujun Yan, Lingxiao Zhao, Mark Heimann, Leman Akoglu, and Danai Koutra. Beyond homophily in graph neural networks: Current limitations and effective designs. In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 7793–7804. Curran Associates, Inc., 2020.
- [101] Marinka Zitnik and Jure Leskovec. Predicting multicellular function through multi-layer tissue networks. *Bioinformatics*, 33(14):i190–i198, 07 2017.

Appendix

Table of Contents

A Related Works	18
A.1 Methods for Alleviating Over-smoothing	18
A.2 Homophily Bias in Evaluation of Techniques for Deep GNN	19
A.3 Methods for Alleviating Over-squashing	20
A.4 Towards a Unified Treatment	21
B Proofs	21
B.1 Expected Propagation Matrix under DropEdge-variants	21
B.2 Sensitivity in L-layer Linear GCNs	23
C Theoretical Extensions	24
C.1 Sensitivity in Nonlinear MPNNs	24
C.2 Test-time Monte-Carlo Dropout	25
D Empirical Sensitivity Analysis	25
D.1 MC-Approximation of Sensitivity in Nonlinear MPNNs	26
D.2 Symmetrically Normalized Propagation Matrix	26
D.3 Upper Bound on Expected Sensitivity	26
D.4 Closer Look at DropSens	27
E Experiments Details	27
E.1 Descriptions of the Datasets	27
E.2 Training Configurations	29
E.3 DropSens Implementation	30
F Supplementary Experiments	31
F.1 Test Accuracy versus DropEdge Probability	31
F.2 Remark on DropNode	32
F.3 Over-squashing or Under-fitting?	32
G Supplementary Experimental Results	33
G.1 Ranking Dropping Methods	33
G.2 Performance of GAT with Dropping Methods	33
G.3 Effect Size in Statistical Tests	33
G.4 Performance of GIN with DropSens	34
G.5 Best-performing Dropping Probabilities	35

A Related Works

A.1 Methods for Alleviating Over-smoothing

A popular choice for reducing over-smoothing in GNNs is to regularize the model. Recall that DropEdge [70] implicitly regularizes the model by adding noise to it (Section 2.2). A similarly regularization effect is observed with the methods discussed in the main text – DropNode [23], DropAgg [43], DropGNN [66], Dropout [79] and DropMessage [21]. Graph Drop Connect (GDC) [37] combines DropEdge and DropMessage together, resulting in a layer-wise sampling scheme that uses a different subgraph for message-aggregation over each feature dimension. These methods successfully addressed the over-smoothing problem, enabling the training of deep GNNs, and performed competitively on several benchmarking datasets.

Table 3: Statistics of node-classification datasets. Homophily measures as defined in [53].

Dataset	Nodes	Edges	Features	Classes	Homophily
Homophilic Networks					
Reddit	232,965	114,615,892	602	41	0.653
OGBN-ArXiv	169,343	1,166,243	128	40	0.416
Coauthor-CS	18,333	163,788	6,805	15	0.755
Coauthor-Physics	34,493	495,924	8,415	5	0.847
Wiki-CS	11,701	216,123	300	10	0.568
Amazon-Computers	13,752	491,722	767	10	0.700
Amazon-Photo	7,650	238,162	745	8	0.772
Heterophilic Networks					
Flickr	89,250	899,756	500	7	0.070
Cornell	183	298	1,703	5	0.031
Texas	183	325	1,703	5	0.001
Wisconsin	251	515	1,703	5	0.094

Another powerful form of implicit regularization is feature normalization, which has proven crucial in enhancing the performance and stability of several types of neural networks [41]. Exploiting the inductive bias in graph-structured data, normalization techniques like PairNorm [95], Differentiable Group Normalization (DGN) [97] and NodeNorm [99] have been proposed to reduce over-smoothing in GNNs. On the other hand, Energetic Graph Neural Networks (EGNNs) [98] explicitly regularize the optimization by constraining the layer-wise Dirichlet energy to a predefined range.

In a different vein, motivated by the success of residual networks (ResNets) [38] in computer vision, [49] proposed the use of residual connections to prevent the smoothing of representations. Residual connections successfully improved the performance of GCN on a range of graph-learning tasks. [13] introduced GCN-II, which uses skip connections from the input to all hidden layers. This layer wise propagation rule has allowed for training of ultra-deep networks – up to 64 layers. Some other architectures, like the Jumping Knowledge Network (JKNet) [86] and the Deep Adaptive GNN (DAGNN) [54], aggregate the representations from *all* layers, $\{z_i^{(\ell)}\}_{\ell=1}^L$, before processing them through a readout layer.

A.2 Homophily Bias in Evaluation of Techniques for Deep GNN

We examine the evaluation protocols commonly used for assessing methods aimed at alleviating over-smoothing in deep GNNs – many of which are also widely adopted for training deep architectures. Notably, we highlight a misalignment between the intended goal of these methods – to improve the trainability of deep GNNs – and their evaluation, which is often restricted to short-range tasks.

For example, DropEdge [70] was evaluated on Cora [58], CiteSeer [31], PubMed [63], and a version of Reddit [36] distinct from the one used in our experiments. The first three exhibit high label homophily (see Table 4) and are known to be better modelled by shallower networks [97]. Reddit also displays strong homophily, as can be seen in Table 3. Similarly, DropNode [23] was evaluated on Cora, CiteSeer, and PubMed; DropAgg [43] on Cora ML, CiteSeer, and OGBN-ArXiv [40], which has moderate homophily; DropMessage [70] was evaluated on Cora, CiteSeer, PubMed, OGBN-ArXiv, and Flickr, with only the latter having low homophily; GDC [37] was evaluated on Cora, Cora ML and CiteSeer.

A similar trend can be observed in the evaluation of feature normalization techniques used to regularize GNNs. PairNorm [95] and DGN [97] were evaluated on Cora, CiteSeer, PubMed, and Coauthor-CS [77]; NodeNorm [99] on Cora, CiteSeer, PubMed, Coauthor-CS, Wiki-CS [59], and Amazon-Photo [77]; and EGNNs [98] on Cora, PubMed, Coauthor-Physics [77], and OGBN-ArXiv – all of these datasets are highly homophilic.

A similar trend is observed in the evaluation of architectural modifications designed to enable deeper GNNs. GCN-II [13] on Cora, CiteSeer, PubMed, and Chameleon; JKNet [86] on Cora, CiteSeer, and Reddit; and DAGNN [54] on Cora, CiteSeer, PubMed, Coauthor-CS, Coauthor-Physics, Amazon-Computers [77], and Amazon-Photo – many of these datasets are highly homophilic as well.

This pattern indicates that an overwhelming proportion of evaluations have been restricted to short-range, homophilic tasks. Such a narrow focus risks overstating the general effectiveness of these methods and masking their potential limitations in long-range scenarios.

A few exceptions stand out. DropGNN [66], which was evaluated on graph-classification from the TUDataset [62], aligning more closely with evaluations of rewiring methods targeting over-squashing [8, 44]. NodeNorm, while primarily evaluated on homophilic datasets, was also tested on three heterophilic graphs: Cornell, Texas, and Wisconsin [68]. GCN-II saw broader evaluation, including on several long-range tasks such as Chameleon, Cornell, Texas, Wisconsin, and the Protein-Protein Interaction (PPI) networks [36]. Lastly, JKNet was also evaluated on the PPI networks.

A.3 Methods for Alleviating Over-squashing

In this section, we will review some of the graph rewiring methods proposed to address the problem of over-squashing. Particularly, we wish to emphasize a commonality among these methods – edge addition is necessary. As a reminder, *graph rewiring* refers to modifying the edge set of a graph by adding and/or removing edges in a systematic manner. In a special case, which includes many of the rewiring techniques we will discuss, the original topology is completely discarded, and only the rewired graph is used for message-passing.

Spatial rewiring methods use the topological relationships between the nodes in order to come up with a rewiring strategy. That is the graph rewiring is guided by the objective of optimizing some chosen topological properties. For instance, [4] introduced a fully-adjacent (FA) layer, wherein messages are passed between all nodes. GNNs using a FA layer in the final message-passing step were shown to outperform the baselines on a variety of long-range tasks, revealing the importance of information exchange between far-off nodes which standard message-passing cannot facilitate. [81] proposed a curvature-based rewiring strategy, called the Stochastic Discrete Ricci Flow (SDRF), which aims to reduce the “bottleneckedness” of a graph by adding suitable edges, while simultaneously removing edges in an effort to preserve the statistical properties of the original topology. [8] proposed the Greedy Total Resistance (GTR) technique, which optimizes the graph’s total resistance by greedily adding edges to achieve the greatest improvement. One concern with graph rewiring methods is that unmoderated densification of the graph, e.g. using a fully connected graph for propagating messages, can result in a loss of the inductive bias the topology provides, potentially leading to over-fitting. Accordingly, [35] propose a Dynamically Rewired (DRew) message-passing framework that gradually *densifies* the graph. Specifically, in a given layer ℓ , node i aggregates messages from its entire ℓ -hop receptive field instead of just the immediate neighbors. This results in an improved communication over long distances while also retaining the inductive bias of the shortest distance between nodes.

Spectral methods, on the other hand, use the spectral properties of the matrices encoding the graph topology, e.g. the adjacency or the Laplacian matrix, to design rewiring algorithms. For example, [5] proposed a differentiable graph rewiring layer based on the Lovász bound [56, Corollary 3.3]. Similarly, [6] introduced the Random Local Edge Flip (RLEF) algorithm, which draws inspiration from the “Flip Markov Chain” [22, 57] – a sequence of such steps can convert a connected graph into an *expander graph* – a sparse graph with good connectivity (in terms of Cheeger’s constant) – with high probability [1, 15, 22, 30, 57], thereby enabling effective information propagation across the graph.

Some other rewiring techniques don’t exactly classify as spatial or spectral methods. For instance, Probabilistically Rewired MPNN (PR-MPNN) [69] learns to probabilistically rewire a graph, effectively mitigating under-reaching as well as over-squashing. Finally, [25] proposed connecting all nodes at most r -hops away, for some $r \in \mathbb{N}$, and introducing positional embeddings to allow for distance-aware aggregation of messages.

A.4 Towards a Unified Treatment

Several studies have shown that an inevitable trade-off exists between the problems of over-smoothing and over-squashing, meaning that optimizing for one will compromise the other. For instance, [64, 81] showed that negatively curved edges create bottlenecks in the graph resulting in over-squashing of information. On the other hand, [64, Proposition 4.3] showed that positively curved edges in a graph contribute towards the over-smoothing problem. To address this trade-off, they proposed Batch Ollivier-Ricci Flow (BORF), which adds new edges adjacent to the negatively curved ones, and simultaneously removes positively curved ones. In a similar vein, [34] demonstrated that the minimum number of message-passing steps required to reach a given level of over-smoothing is inversely related to the Cheeger’s constant, h_G . This again implies an inverse relationship between over-smoothing and over-squashing. To effectively alleviate the two issues together, they proposed the Stochastic Jost and Liu Curvature Rewiring (SJLR) algorithm, which adds edges that result in high improvement in the curvature of existing edges, while simultaneously removing those that have low curvature.

Despite the well-established trade-off between over-smoothing and over-squashing, some works have successfully tackled them together despite *only* adding or removing edges. One such work is [44], which proposed a rewiring algorithm that adds edges to the graph but does not remove any. The First-order Spectral Rewiring (FoSR) algorithm computes, as the name suggests, a first-order approximation to the spectral gap of the symmetric Laplacian matrix ($\mathbf{L}^{\text{sym}} = \mathbf{I}_N - (\mathbf{D}^\dagger)^{1/2} \mathbf{A} (\mathbf{D}^\dagger)^{1/2}$), and adds edges with the aim of maximizing it. Since the spectral gap directly relates to Cheeger’s constant – a measure of bottleneck-edness in the graph – through Cheeger’s inequality [2, 3, 78], this directly decreases the over-squashing levels. Moreover, [44, Figure 5] empirically demonstrated that addition of (up to a small number of) edges selected by FoSR can lower the Dirichlet energy of the representations, suggesting the method’s potential to simultaneously tackle over-smoothing. Taking a somewhat opposite approach, [55] adapted DropEdge to *remove* negatively curved edges sampled from a distribution proportional to edge curvatures. Their method, called CurvDrop, directly reduces over-squashing and, as a side benefit of operating on a sparser subgraph, also mitigates over-smoothing.

B Proofs

B.1 Expected Propagation Matrix under DropEdge-variants

Lemma. *When using DropEdge, the expected propagation matrix is given as:*

$$\begin{aligned}\mathbb{E}_{\text{DE}} \left[\hat{\mathbf{A}}_{ii}^{(1)} \right] &= \frac{1 - q^{d_i+1}}{(1 - q)(d_i + 1)} \\ \mathbb{E}_{\text{DE}} \left[\hat{\mathbf{A}}_{ij}^{(1)} \right] &= \frac{1}{d_i} \left(1 - \frac{1 - q^{d_i+1}}{(1 - q)(d_i + 1)} \right)\end{aligned}$$

where $(j \rightarrow i) \in \mathcal{E}$; $\hat{\mathbf{P}}_{ij} = 0$ otherwise.

Proof. Recall that under DropEdge, a self-loop is added to the graph *after* the edges are dropped, and then the normalization is performed. In other words, the self-loop is never dropped. Therefore, given the i.i.d. masks, $m_1, \dots, m_{d_i} \sim \text{Bern}(1 - q)$, on incoming edges to node i , the total number of messages is given by

$$1 + \sum_{k=1}^{d_i} m_k = 1 + M_i \tag{B.1}$$

where $M_i \sim \text{Binom}(d_i, 1 - q)$. Under asymmetric normalization (see Section 2.1), the expected weight of the message along the self-loop is computed as follows:

$$\mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ii}^{(1)}] = \mathbb{E}_{m_1, \dots, m_{d_i}} \left[\frac{1}{1 + \sum_{k=1}^{d_i} m_k} \right] \quad (\text{B.2})$$

$$= \mathbb{E}_{M_i} \left[\frac{1}{1 + M_i} \right] \quad (\text{B.3})$$

$$= \sum_{k=0}^{d_i} \binom{d_i}{k} (1 - q)^k (q)^{d_i - k} \left(\frac{1}{1 + k} \right) \quad (\text{B.4})$$

$$= \frac{1}{(1 - q)(d_i + 1)} \sum_{k=0}^{d_i} \binom{d_i + 1}{k + 1} (1 - q)^{k+1} (q)^{d_i - k} \quad (\text{B.5})$$

$$= \frac{1}{(1 - q)(d_i + 1)} \sum_{k=1}^{d_i + 1} \binom{d_i + 1}{k} (1 - q)^k (q)^{d_i + 1 - k} \quad (\text{B.6})$$

$$= \frac{1 - q^{d_i + 1}}{(1 - q)(d_i + 1)} \quad (\text{B.7})$$

Similarly, if the Bernoulli mask corresponding to $j \rightarrow i$ is 1, then the total number of incoming messages to node i is given by

$$2 + \sum_{k=1}^{d_i - 1} m_k$$

including one self-loop, which is never dropped, as noted earlier. On the other hand, the weight of the edge is simply 0 if the corresponding Bernoulli mask is 0. Using the Law of Total Expectation, the expected weight of the edge $j \rightarrow i$ can be computed as follows:

$$\mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ij}^{(1)}] = q \cdot 0 + (1 - q) \mathbb{E}_{m_1, \dots, m_{d_i - 1}} \left[\frac{1}{2 + \sum_{k=1}^{d_i - 1} m_k} \right] \quad (\text{B.8})$$

$$= (1 - q) \sum_{k=0}^{d_i - 1} \binom{d_i - 1}{k} (1 - q)^k (q)^{d_i - 1 - k} \left(\frac{1}{2 + k} \right) \quad (\text{B.9})$$

$$= \sum_{k=0}^{d_i - 1} \frac{(d_i - 1)!}{(k + 2)!(d_i - 1 - k)!} (1 - q)^{k+1} (q)^{d_i - 1 - k} (k + 1) \quad (\text{B.10})$$

$$= \sum_{k=2}^{d_i + 1} \frac{(d_i - 1)!}{(k)!(d_i + 1 - k)!} (1 - q)^{k-1} (q)^{d_i + 1 - k} (k - 1) \quad (\text{B.11})$$

$$= \frac{1}{d_i(d_i + 1)(1 - q)} \sum_{k=2}^{d_i + 1} \binom{d_i + 1}{k} (1 - q)^k (q)^{d_i + 1 - k} (k - 1) \quad (\text{B.12})$$

$$= \frac{1}{d_i(d_i + 1)(1 - q)} [(d_i + 1)(1 - q) - 1 + q^{d_i + 1}] \quad (\text{B.13})$$

$$= \frac{1}{d_i} \left(1 - \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ii}^{(1)}] \right) \quad (\text{B.14})$$

□

Analysis of DropEdge-variants. We will similarly derive the expected propagation matrix for other random edge-dropping algorithms. First off, DropNode [23] samples nodes and drops corresponding columns from the aggregation matrix directly, followed by rescaling of its entries:

$$\mathbb{E}_{\text{DN}} \left[\frac{1}{1 - q} \hat{\mathbf{A}} \right] = \frac{1}{1 - q} \times (1 - q) \hat{\mathbf{A}} = \hat{\mathbf{A}} \quad (\text{B.15})$$

That is, the expected propagation matrix is the same as in a NoDrop model ($q = 0$).

Nodes sampled by DropAgg [43] don't aggregate messages. Therefore, if $\hat{\mathbf{A}} = \hat{\mathbf{A}}^{\text{asym}}$, then the expected propagation matrix is given by

$$\begin{aligned}\mathbb{E}_{\text{DA}} [\hat{\mathbf{A}}_{ii}] &= q + \frac{1-q}{d_i+1} = \frac{1+d_i q}{d_i+1} > \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ii}] \\ \mathbb{E}_{\text{DA}} [\hat{\mathbf{A}}_{ij}] &= \frac{1}{d_i} \left(1 - \mathbb{E}_{\text{DA}} [\hat{\mathbf{A}}_{ii}]\right) < \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ij}]\end{aligned}\quad (\text{B.16})$$

Finally, DropGNN [66] samples nodes which neither propagate nor aggregate messages. From any node's perspective, if it is not sampled, then its aggregation weights are computed as for DropEdge:

$$\begin{aligned}\mathbb{E}_{\text{DG}} [\hat{\mathbf{A}}_{ii}] &= q + (1-q) \mathbb{E}_{\text{DE}} [\hat{\mathbf{A}}_{ii}] = q + \frac{1-q^{d_i+1}}{d_i+1} > \mathbb{E}_{\text{DA}} [\hat{\mathbf{A}}_{ii}] \\ \mathbb{E}_{\text{DG}} [\hat{\mathbf{A}}_{ij}] &= \frac{1}{d_i} \left(1 - \mathbb{E}_{\text{DG}} [\hat{\mathbf{A}}_{ii}]\right) < \mathbb{E}_{\text{DA}} [\hat{\mathbf{A}}_{ij}]\end{aligned}\quad (\text{B.17})$$

B.2 Sensitivity in L-layer Linear GCNs

Theorem. In an L-layer linear GCN with $\hat{\mathbf{A}} = \hat{\mathbf{A}}^{\text{asym}}$, using DropEdge, DropAgg or DropGNN decreases the sensitivity of a node $i \in \mathcal{V}$ to another node $j \in \mathbb{S}^{(L)}(i)$, thereby reducing its effective receptive field.

$$\mathbb{E}_{\dots} \left[\left(\hat{\mathbf{A}}^L \right)_{ij} \right] = \sum_{(u_0, \dots, u_L) \in \text{Paths}(j \rightarrow i)} \prod_{\ell=1}^L \mathbb{E}_{\dots} [\hat{\mathbf{A}}_{u_\ell u_{\ell-1}}] < \mathbb{E}_{\text{ND}} \left[\left(\hat{\mathbf{A}}^L \right)_{ij} \right] \quad (\text{B.18})$$

where ND refers to a NoDrop model ($q = 0$), the placeholder $\dots$ can be replaced with one of the edge-dropping methods DE, DA or DG, and the corresponding entries of $\mathbb{E}_{\dots}[\hat{\mathbf{A}}]$ can be plugged in from Equation 3.3, Equation B.16 and Equation B.17, respectively. Moreover, the sensitivity monotonically decreases as the dropping probability is increased.

Proof. Recall that $\dot{\mathbf{P}}$ can be viewed as the transition matrix of a non-uniform random walk, such that $\dot{\mathbf{P}}_{uv} = \mathbb{P}(u \rightarrow v)$. Intuitively, since there is no self-loop on any given L-length path connecting nodes i and j (which are assumed to be L-hops away), the probability of each transition on any path connecting these nodes is reduced. Therefore, so is the total probability of transitioning from i to j in exactly L hops.

More formally, denote the set of paths connecting the two nodes by

$$\text{Paths}(j \rightarrow i) = \{(u_0, \dots, u_L) : u_0 = j; u_L = i; (u_{\ell-1} \rightarrow u_\ell) \in \mathcal{E}, \forall \ell \in [L]\} \quad (\text{B.19})$$

The (i, j) -entry in the propagation matrix is given by

$$\left(\dot{\mathbf{P}}^L \right)_{ij} = \sum_{(u_0, \dots, u_L) \in \text{Paths}(j \rightarrow i)} \prod_{\ell=1}^L \dot{\mathbf{P}}_{u_\ell u_{\ell-1}} \quad (\text{B.20})$$

Since there is no self-loop on any of these paths,

$$\left(\dot{\mathbf{P}}^L \right)_{ij} = \sum_{(u_0, \dots, u_L) \in \text{Paths}(j \rightarrow i)} \prod_{\ell=1}^L \frac{1}{d_{u_\ell}} \left(1 - \frac{1 - q^{d_{u_\ell}+1}}{(1-q)(d_{u_\ell}+1)} \right) \quad (\text{B.21})$$

$$< \sum_{(u_0, \dots, u_L) \in \text{Paths}(j \rightarrow i)} \prod_{\ell=1}^L \left(\frac{1}{d_{u_\ell} + 1} \right) \quad (\text{B.22})$$

The right hand side of the inequality is the (i, j) -entry in the L^{th} power of the propagation matrix of a NoDrop model. From Equation B.16 and Equation B.17, we know that Equation B.22 is true for

DropAgg and DropGNN as well. We conclude the first part of the proof using Equation 3.2 – the sensitivity of node i to node j is proportional to $(\dot{\mathbf{P}}^L)_{ij}$.

Next, we recall the geometric series for any q :

$$1 + q + \dots + q^d = \frac{1 - q^{d+1}}{1 - q} \quad (\text{B.23})$$

Each of the terms on the right are increasing in q , hence, all the $\dot{\mathbf{P}}_{u_\ell u_{\ell-1}}$ factors are decreasing in q . Similarly, $\mathbb{E}_{\text{DA}}[\hat{\mathbf{A}}_{ij}]$ and $\mathbb{E}_{\text{DG}}[\hat{\mathbf{A}}_{ij}]$ decrease with increasing q . Using these results with Equation B.20, we conclude the second part of the theorem. $\square$

C Theoretical Extensions

C.1 Sensitivity in Nonlinear MPNNs

While linear networks are useful in simplifying the theoretical analysis, they are often not practical. In this subsection, we will consider the upper bounds on sensitivity established in previous works, and extend them to the DropEdge setting.

ReLU GCNs. [86] considered the case of ReLU nonlinearity, so that the update rule is $\mathbf{Z}^{(\ell)} = \text{ReLU}(\hat{\mathbf{A}}\mathbf{Z}^{(\ell-1)}\mathbf{W}^{(\ell)})$. Additionally, it makes the simplifying assumption that each path in the computational graph is *active* with a fixed probability, ρ [45, Assumption A1p-m]. Accordingly, the sensitivity (in expectation) between any two nodes is given as

$$\left\| \mathbb{E}_{\text{ReLU}} \left[\frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right] \right\|_1 = \left[\rho \left\| \prod_{\ell=1}^L \mathbf{W}^{(\ell)} \right\|_1 \right] (\hat{\mathbf{A}}^L)_{ij} = \zeta_1^{(L)} (\hat{\mathbf{A}}^L)_{ij} \quad (\text{C.1})$$

where $\zeta_1^{(L)}$ depends only on the depth L , and is independent of the choice of nodes $i, j \in \mathcal{V}$. Taking an expectation w.r.t. the random edge masks, we get

$$\mathbb{E}_{\mathbf{M}^{(1)}, \dots, \mathbf{M}^{(L)}} \left[\left\| \mathbb{E}_{\text{ReLU}} \left[\frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right] \right\|_1 \right] = \zeta_1^{(L)} \left(\mathbb{E} \left[\prod_{\ell=1}^L \hat{\mathbf{A}}^{(\ell)} \right] \right)_{ij} = \zeta_1^{(L)} \left(\mathbb{E} [\hat{\mathbf{A}}]^L \right)_{ij} \quad (\text{C.2})$$

Using Theorem 3.1, we conclude that in a ReLU-GCN, DropEdge, DropAgg and DropGNN will reduce the expected sensitivity between nodes L -hops away. Empirical observations in Figures 1a and 4 suggest that we may expect an increase in sensitivity to neighboring nodes, but a significant decrease in sensitivity to those farther away.

Source-only Message Functions. [8, Lemma 3.2] considers MPNNs with aggregation functions of the form

$$\text{Agg}^{(\ell)} \left(\mathbf{z}_i^{(\ell-1)}, \left\{ \mathbf{z}_j^{(\ell-1)} : j \in \mathbb{S}^{(1)}(i) \right\} \right) = \sum_{j \in \mathbb{B}^{(1)}(i)} \hat{\mathbf{A}}_{ij} \text{Msg}^{(\ell)} \left(\mathbf{z}_j^{(\ell-1)} \right) \quad (\text{C.3})$$

and Upd and Msg functions with bounded gradients. In this case, the sensitivity between two nodes $i, j \in \mathcal{V}$ can be bounded as

$$\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \leq \zeta_2^{(L)} \left(\sum_{\ell=0}^L \hat{\mathbf{A}}^\ell \right)_{ij} \quad (\text{C.4})$$

As before, we can use the independence of edge masks to get an upper bound on the expected sensitivity:

$$\mathbb{E}_{\mathbf{M}^{(1)}, \dots, \mathbf{M}^{(L)}} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] \leq \zeta_2^{(L)} \left(\mathbb{E} \left[I_N + \sum_{\ell=1}^L \prod_{k=1}^{\ell} \hat{\mathbf{A}}^{(k)} \right] \right)_{ij} = \zeta_2^{(L)} \left(\sum_{\ell=0}^L \mathbb{E} [\hat{\mathbf{A}}]^\ell \right)_{ij} \quad (\text{C.5})$$

Figure 5 shows the plot of the entries of $\sum_{\ell=0}^6 \dot{P}^\ell$ (i.e. for DropEdge), as in the upper bound above, with $\hat{A} = \hat{A}^{\text{asym}}$. We observe that the sensitivity between nearby nodes marginally increases, while that between distant nodes notably decreases (similar to Figure 1a), suggesting significant over-squashing. Similar observations can be made with $\hat{A} = \hat{A}^{\text{sym}}$, and for other DropEdge-variants.

Source-and-Target Message Functions. [81, Lemma 1] showed that if the aggregation function is instead given by

$$\text{Agg}^{(\ell)} \left(\mathbf{z}_i^{(\ell-1)}, \left\{ \mathbf{z}_j^{(\ell-1)} : j \in \mathbb{S}^{(1)}(i) \right\} \right) = \sum_{j \in \mathbb{B}^{(1)}(i)} \hat{A}_{ij} \text{Msg}^{(\ell)} \left(\mathbf{z}_i^{(\ell-1)}, \mathbf{z}_j^{(\ell-1)} \right) \quad (\text{C.6})$$

then the sensitivity between nodes $i \in \mathcal{V}$ and $j \in \mathbb{S}^{(L)}(i)$ can be bounded as

$$\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \leq \zeta_3^{(L)} \left(\hat{A}^L \right)_{ij} \quad (\text{C.7})$$

With random edge-dropping, this bound can be adapted as follows:

$$\mathbb{E}_{\mathbf{M}^{(1)}, \dots, \mathbf{M}^{(L)}} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] \leq \zeta_3^{(L)} \left(\mathbb{E} \left[\hat{A} \right]^L \right)_{ij} \quad (\text{C.8})$$

which is similar to Equation C.2, only with a different proportionality constant, that is anyway independent of the choice of nodes. Here, again, we invoke Theorem 3.1 to conclude that $(\mathbb{E}[\hat{A}^L])_{ij}$ decreases monotonically with increasing DropEdge probability q . This implies that, in a non-linear MPNN with $\hat{A} = \hat{A}^{\text{asym}}$, DropEdge lowers the sensitivity bound given above. Empirical results in Figure 4 support the same conclusion for $\hat{A} = \hat{A}^{\text{sym}}$.

C.2 Test-time Monte-Carlo Dropout

Up until now, we have focused on the expected sensitivity of the stochastic representations in models using DropEdge-variants. This corresponds to their training-time behavior, wherein the activations are random. At test-time, the standard practice is to turn these methods off by setting $q = 0$. However, this raises the over-smoothing levels back up [88]. Another way of making predictions is to perform multiple stochastic forward passes, as during training, and then averaging the model outputs. This is similar to Monte-Carlo Dropout, which is an efficient way of ensemble averaging in MLPs [27], CNNs [26] and RNNs [28]. In addition to alleviating over-smoothing, this approach also outperforms the standard implementation in practical settings [88]. We can study the effect of random edge-dropping in this setting by examining the sensitivity of the *expected representations*:

$$\left\| \frac{\partial}{\partial \mathbf{x}_j} \mathbb{E} \left[\mathbf{z}_i^{(L)} \right] \right\|_1 = \left\| \mathbb{E} \left[\frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right] \right\|_1 \quad (\text{C.9})$$

In linear models, the order of the two operations – expectation and 1-norm – is irrelevant:

$$\left\| \mathbb{E} \left[\frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right] \right\|_1 = \left\| \mathbb{E} \left[\left(\hat{A}^L \right)_{ij} \mathbf{W} \right] \right\|_1 = \mathbb{E} \left[\left(\hat{A}^L \right)_{ij} \|\mathbf{W}\|_1 \right] = \mathbb{E} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] \quad (\text{C.10})$$

In general, the two quantities can be related using the convexity of norms and Jensen's inequality:

$$\left\| \frac{\partial}{\partial \mathbf{x}_j} \mathbb{E} \left[\mathbf{z}_i^{(L)} \right] \right\|_1 \leq \mathbb{E} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] \leq \dots \quad (\text{C.11})$$

Therefore, the upper bound results in Appendix C.1 trivially extend to the MC-averaged representations. Although tighter bounds may be derived for this setting, we leave that for future works.

D Empirical Sensitivity Analysis

In this section, we elaborate on the setup used for the empirical sensitivity analysis in Section 3. We also present some supplemental figures demonstrating the negative effects of random edge-dropping, particularly focusing on scenarios not covered by the theory.

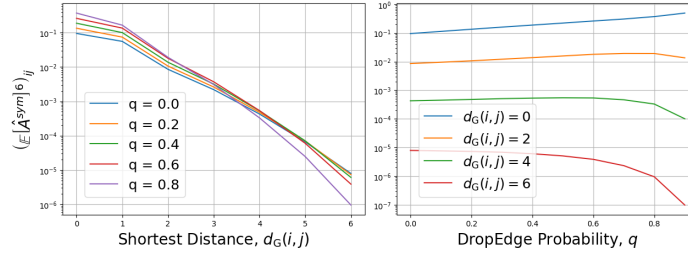


Figure 4: Entries of $\ddot{P}^6$, averaged after binning node-pairs by their shortest distance.

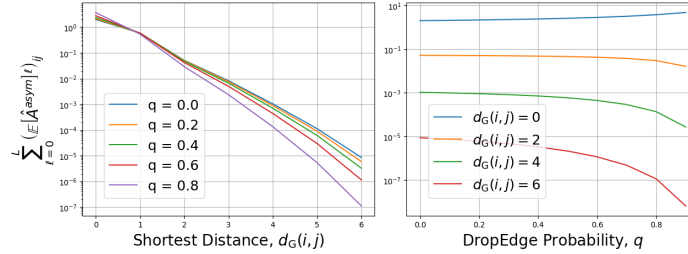


Figure 5: Entries of $\sum_{\ell=0}^6 \ddot{P}^\ell$, averaged after binning node-pairs by their shortest distance.

D.1 MC-Approximation of Sensitivity in Nonlinear MPNNs

Given a target node from the Cora dataset [58], we computed the sensitivity of its representation to source nodes up to $L = 6$ hops away in ReLU-GCNs of width 32. The raw sensitivities were normalized to obtain *influence scores* [86]. This was repeated for 25 target nodes, and 25 model–dropout samples were used for each of them. The source nodes were binned by the shortest distance from the corresponding target node, and the influence scores were averaged over each bin to obtain an *average influence* from nodes ℓ -hops away.

D.2 Symmetrically Normalized Propagation Matrix

The results in Section 3 correspond to the use of $\hat{A} = \hat{A}^{\text{asym}}$ for aggregating messages – in each message passing step, only the in-degree of node i is used to compute the aggregation weights of the incoming messages. In practice, however, it is more common to use the symmetrically normalized propagation matrix, $\hat{A} = \hat{A}^{\text{sym}}$, which ensures that nodes with high out-degree do not dominate the information flow in the graph [47]. As in Equation 3.2, we are looking for

$$\ddot{P}^L := \mathbb{E}_{M^{(1)}, \dots, M^{(L)}} \left[\prod_{\ell=1}^L \hat{A}^{(\ell)} \right] \quad (\text{D.1})$$

where $\ddot{P} := \mathbb{E}_{\text{DE}}[\hat{A}^{\text{sym}}]$. While $\ddot{P}$ is analytically intractable, we can approximate it using Monte-Carlo sampling. Accordingly, we use the Cora dataset, and sample 20 DropEdge masks to compute an approximation of $\ddot{P}$, and plot out the entries of $\ddot{P}^L$, as we did for $\dot{P}^L$ in Figure 1a. The results are presented in Figure 4, which shows that while the sensitivity between nodes up to 3 hops away is increased, that between nodes farther off is significantly reduced, same as in Figure 1a.

D.3 Upper Bound on Expected Sensitivity

[8] showed that the sensitivity between any two nodes in a graph can be bounded using the sum of the powers of the propagation matrix. In Appendix C.1, we extended this bound to random edge-dropping methods with independent edge masks sampled in each layer:

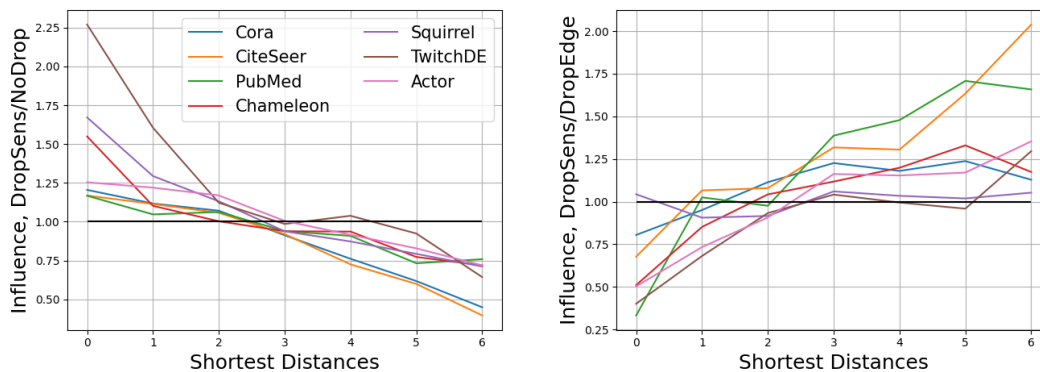


Figure 6: Average sensitivity under DropSens, normalized by average sensitivity under NoDrop (*left*) and DropEdge (*right*). We note that DropSens is more sensitive to nearby nodes ($d \leq 2$) than NoDrop, and more sensitive to far off nodes ($d \geq 3$) than DropEdge.

$$\mathbb{E}_{\mathbf{M}^{(1)}, \dots, \mathbf{M}^{(L)}} \left[\left\| \frac{\partial \mathbf{z}_i^{(L)}}{\partial \mathbf{x}_j} \right\|_1 \right] \leq \zeta_3^{(L)} \left(\sum_{\ell=0}^L \mathbb{E} [\hat{\mathbf{A}}^\ell] \right)_{ij}$$

Although this bound does not have a closed form, we can again use the Cora network to study its entries. We plot the entries of $\sum_{\ell=0}^L \dot{\mathbf{P}}^\ell$, corresponding to DropEdge, against the shortest distance between node-pairs. The results are presented in Figure 5. We observe an exponential decrease in the sensitivity bound as the distance between nodes increases, suggesting that DropEdge is not suitable for capturing LRIs.

D.4 Closer Look at DropSens

Lastly, we evaluate DropSens on node-classification tasks, comparing its effects with NoDrop and DropEdge. In Figure 6, we compare the sensitivity profiles of these methods. It is clear to observe that DropSens allows node representations to be more sensitive to distant nodes' ($d \geq 3$) features, than DropEdge, thereby improving performance on heterophilic datasets. What is interesting is that for homophilic datasets, DropSens is more sensitive to nearby nodes ($d \geq 1$) than DropEdge, allowing the model to effectively capture short-range interactions. Although it is poorer at capturing short-range interactions for heterophilic datasets, DropSens improves over NoDrop on that front, similar to DropEdge.

This set of experiments suggests that DropSens strikes a perfect balance between NoDrop and DropEdge:

1. improves sensitivity to nearby nodes ($d \leq 2$), over NoDrop,
2. improves sensitivity to distant nodes ($d \geq 3$), over DropEdge, and
3. mitigates over-fitting (by stochastic regularization) and over-smoothing (by reducing message-passing), just as DropEdge.

E Experiments Details

In this section, we expand on the details of the experiments in Section 5. All experiments were run on a server equipped with an Intel(R) Xeon(R) E5-2620 v3 CPU, 62 GB of RAM, $4 \times$ NVIDIA GeForce GTX TITAN X GPU (12 GB VRAM each), and CUDA version 12.4.

E.1 Descriptions of the Datasets

Synthetic Datasets. The SyntheticZINC dataset [33], as the name suggests, is a synthetic dataset derived from the ZINC chemical dataset [42], with the dataset size constrained to 12K molecular

Table 4: Statistics of node-classification datasets. Homophily measures from [53].

Dataset	Nodes	Edges	Features	Classes	Homophily
Homophilic Networks					
Cora	2,708	10,556	1,433	7	0.766
CiteSeer	3,327	9,104	3,703	6	0.627
PubMed	19,717	88,648	500	3	0.664
Heterophilic Networks					
Chameleon	2,277	36,051	2,325	5	0.062
Squirrel	5,201	216,933	2,089	5	0.025
Actor	7,600	29,926	931	5	0.011
TwitchDE	9,498	306,276	128	2	0.142

Table 5: Statistics of graph-classification datasets.

Dataset	Graphs	Nodes	Edge	Features	Classes
Mutag	188	17.9	39.6	7	2
Proteins	1,113	39.1	145.6	3	2
Enzymes	600	32.6	124.3	3	6
Collab	5,000	74.5	4914.4	0	3
IMDb	1,000	19.8	193.1	0	2
Reddit	2,000	429.6	995.5	0	2

graphs [20]. Specifically, given a molecular graph $\mathcal{G}$, we set all its nodes' features to 0, except for two nodes, i and $j \neq i$, whose features are sampled as $x_i, x_j \in \mathcal{U}(0, 1)$. The graph-level target is computed as $y = \tanh(x_i + x_j)$, i.e. learning the task requires a non-linear mixing between the features of nodes i and j . These nodes are chosen to induce the desired level of underlying mixing – given $\alpha \in [0, 1]$, the node-pair (i, j) is chosen such that the commute time [10] between them is the α^{th} quantile of the distribution of commute times over $\mathcal{G}$. We analyze the effect of underlying mixing on model performance by varying α as 0.1, 0.2, $\dots$, 1.0. The MPNN is chosen to be an L-layer GCN with a MAX-pooling readout, which encourages the model to learn the mixing by effectively passing messages [33, Theorem 3.2]. The model depth is set at $L = \max_{\mathcal{G}} \lceil \text{diam}(\mathcal{G})/2 \rceil = 11$ to ensure that the GCN does not suffer from under-reaching [4, 7].

Node-classification Tasks. Cora [58], CiteSeer [31] and PubMed [63] are citation networks – their nodes represent scientific publications and an edge between two nodes indicates that one of them has cited the other. The features of each publication are represented by a binary vector, where each index indicates whether a specific word from a dictionary is present or absent. Several studies have showed that these datasets have high homophily in node labels [53, 100] and that they are modelled much better by shallower networks than by deeper ones [97]. Chameleon and Squirrel [71] are networks of English Wikipedia web pages on the respective topics, and the edges between web pages indicate links between them. The task is to predict the average-monthly traffic on each of the web pages. The Actor dataset is induced from a larger film-director-actor-writer network [68]. It is a network of film actors, with edges between those that occur on the same Wikipedia page, and node features are binary vectors denoting the presence of specific keywords in the corresponding Wikipedia entries. The task is to classify actors into five categories based on the content of their Wikipedia pages. Finally, TwitchDE [71] is a network of Twitch users in Germany, with the edges between them representing their mutual follower relationships. The node features are embeddings of the games played by the users, and the task is to predict whether the users use explicit language.

Graph-classification Tasks. Following [8, 44], we use the graph-classification datasets introduced in [62], which were hypothesized to be long-range tasks. On one hand we have the molecular datasets Mutag, Proteins and Enzymes, and on another we have the social networks Reddit-Binary, IMDb-Binary and Collab. Their statistics are presented in Table 5.

Mutag [17] consists of nitroaromatic compounds, with the objective of predicting their mutagenicity, i.e. the ability to cause genetic mutations, in cells in *Salmonella typhimurium*. Each compound is represented as a graph, where nodes correspond to atoms, represented by their type using one-hot encoding, and edges denote the bonds between them. Proteins [9, 19] is a collection of proteins classified as enzymes or not. The molecules are represented as a graph of amino acids, with edges between those separated up to 6Å apart. The Enzymes [9, 75] dataset is represented similarly, and the task is to classify the proteins into one of 6 Enzyme Commission (EC) numbers – a system to classify enzymes based on the reactions they catalyze.

Collab [89] is a scientific collaboration dataset where each graph represents a researcher’s ego network. Nodes correspond to researchers and their collaborators, with edges indicating co-authorship. Each network is labeled based on the researcher’s field, which can be High Energy Physics, Condensed Matter Physics, or Astrophysics. IMDb-Binary [89] is a movie collaboration dataset containing the ego-networks of 1,000 actors and actresses from IMDB. In each graph, nodes represent actors, with edges connecting those who have co-starred in the same film. Each graph is derived from either the Action genre or Romance. Finally, Reddit-Binary [89] comprises graphs representing online discussions on Reddit, where nodes correspond to users and edges indicate interactions through comment responses. Each graph is labeled based on whether it originates from a Q&A or discussion-based subreddit.

For the molecular datasets, we use the node features supplied by PyG [24], but since they are unavailable for the social networks, we set scalar features $x_i = (1)$ for all nodes in these datasets, following [44].

E.2 Training Configurations

Dataset Splits. For the SyntheticZINC task, we use the train-val-test splits provided in PyG. For the homophilic (citation) networks, we use the ‘full’ split [12], as provided in PyG, and for the heterophilic networks, we randomly sample 60% of the nodes for training, 16% for validation, and 24% for testing. On the other hand, for the graph classification tasks, we sample 80% of the graphs for training, and 10% each for validation and testing, following [8, 44].

Model Architecture. We standardize most of the hyperparameters across all experiments to isolate the effect of random dropping. Specifically, we use symmetric normalization of the adjacency matrix to compute the edge weights for GCN, and we set the number of attentions heads for GAT to 2 in order to keep the computational load manageable, while at the same time harnessing the expressiveness of the multi-headed self-attention mechanism. For the SyntheticZINC dataset, we fix the size of the hidden representations at 16, while we fix them to 64 for all the real-world datasets. In all settings, a linear transformation is applied to the node features before message-passing. Afterwards, a bias term is added and then the ReLU nonlinearity is applied. Finally, a linear readout layer is used to compute the regressand (for regression tasks) or logits (for classification).

Dropping Probability. For the synthetic datasets, we experiment with a NoDrop baseline, and DropEdge, Dropout and DropMessage, each with $q = 0.2$ and $q = 0.5$. For the real-world datasets, the dropping probabilities are varied as $q = 0.1, 0.2, \dots, 0.9$, so as to reliably find the best performing configuration. We adopt the common practice of turning the dropping methods off at test-time ($q = 0$), isolating the effects on optimization and generalization, which our theory does not address.

DropSens Configurations. For DropSens, we use 4 possible values for proportion of information preserved over corss-edges, $c = 0.5, 0.8, 0.9, 0.95$. Since the dropping probability, q_i , increases with the in-degree of the target node, d_i , the proportion of all edges dropped could become very high, especially with small c . Therefore, we clip the value of q_i by 4 possible choices, $q_i \leq q_{\max} \in \{0.2, 0.3, 0.5, 0.8\}$. We exclude the following configurations: $(c, q_{\max}) \in \{(0.5, 0.2), (0.5, 0.3), (0.5, 0.5), (0.8, 0.2)\}$, since they use the same dropping probability ($= q_{\max}$) for each edge, and are therefore equivalent to DropEdge. In summary, we test with a total of 12 configurations for DropSens to find the best one for each task.

Optimization Algorithm. The models are trained using the Adam optimizer [46]. On the SyntheticZINC dataset, the models are trained with a learning rate of 2×10^{-3} and a weight decay of 1×10^{-4} , for a total of 200 epochs. On the real-world datasets, we use a learning rate of 1×10^{-3} and no weight decay, following [8, 44]. Here, we cap the maximum number of epochs at 300. In

both cases, the learning rate is reduced by a factor of 1×10^{-1} if the validation loss fails to improve beyond a relative threshold of 1×10^{-4} for 10 epochs, again following [8, 44].

Number of Independent Runs. We perform only 10 independent runs on the SyntheticZINC dataset due to its consistently low variance in performance, as also observed by [33]. For real-world datasets, we conduct 20 runs to identify the best-performing dropping configurations. We then perform a one-sided t-test to assess whether dropout improves performance, using a 90% confidence-level ($\alpha = 0.1$) and targeting a statistical power of 0.9 ($\beta = 0.1$). Under the assumption that the dropping method offers superior performance, detecting a medium effect size of 0.5 [14] requires approximately 53 samples per group according to standard power analysis. We round this to 50, and accordingly perform 30 additional runs for the final comparison of the best-performing dropping configuration with the NoDrop baseline.

E.3 DropSens Implementation

```
import warnings
import sympy
from sympy.abc import x as q
import torch
from torch_geometric.utils import degree, contains_self_loops

def drop_sens(
    edge_index: torch.Tensor,
    c: float,
    max_drop_prob: float = None
):
    if max_drop_prob is None:
        max_drop_prob = 1.

    # Assuming edge index does not have self loops
    if contains_self_loops(edge_index):
        warnings.warn("Degree computation in DropSens assumes absence"
            " of self-loops, but the edge_index passed contains some.")
    degrees = degree(edge_index[1]).int() # Node index -> in-degree

    ds = torch.unique(degrees).tolist() # Sorted array
    mapper = torch.nan * torch.ones(ds[-1]+1)
    mapper[ds] = max_drop_prob # Node in-degree -> dropping prob

    for d_i in ds:
        q_i = float(sympy.N(sympy.real_roots(
            (1-c)*d_i*(1-q) - q + q**(d_i+1))[-2] # Following Equation 4.1
        )) if d_i > 0 else 0.
        if q_i > max_drop_prob:
            # Because q monotononic wrt d, and ds is sorted
            break
        mapper[d_i] = q_i

    in_degrees = degrees[edge_index[1]] # Edge index -> in-degree of target node
    qs = mapper[in_degrees] # Edge index -> dropping probability
    edge_mask = qs <= torch.rand(edge_index.size(1))
    edge_index = edge_index[:, edge_mask]

    return edge_index, edge_mask
```

Listing 1: DropSens Implementation

In Listing 1, we present the DropSens implementation used in our experiments, relying mainly on SymPy [60].

Unfortunately, computing the roots of Equation 4.1 becomes slow when the in-degree d_i is large – a common scenario in large networks. This issue is especially pronounced when the proportion of information preserved c is large, as the dropping threshold is only met at a higher value of d_i . To address this computational challenge, we propose an approximation:

$$1 - c = \frac{q_i - q_i^{d_i+1}}{d_i(1 - q_i)} \approx \frac{q_i}{d_i(1 - q_i)} \implies q_i \approx \frac{(1 - c) d_i}{1 + (1 - c) d_i} \quad (\text{E.1})$$

This approximation becomes increasingly accurate as c increases – since more information needs to be preserved, q_i needs to be small, and hence, $q_i^{d_i+1} \rightarrow 0$.

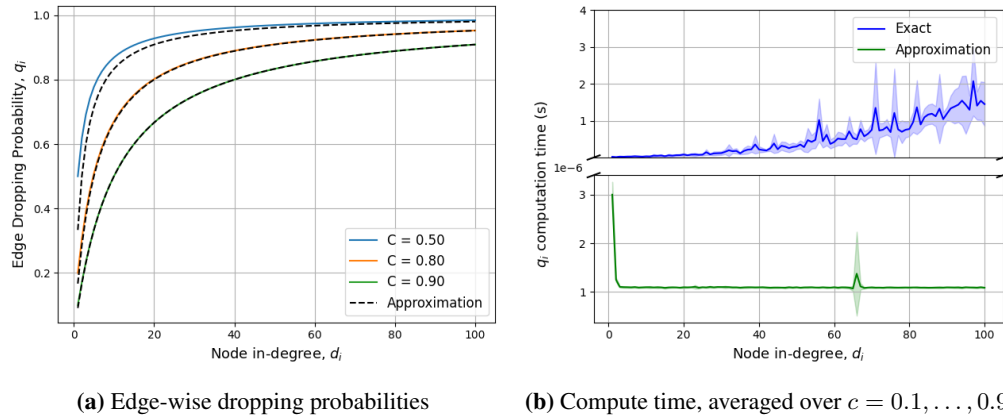


Figure 7: Edge-wise dropping probabilities under DropSens for varying values of c , along with the approximation as in Equation E.1. Compute times are measured using only 1 execution thread.

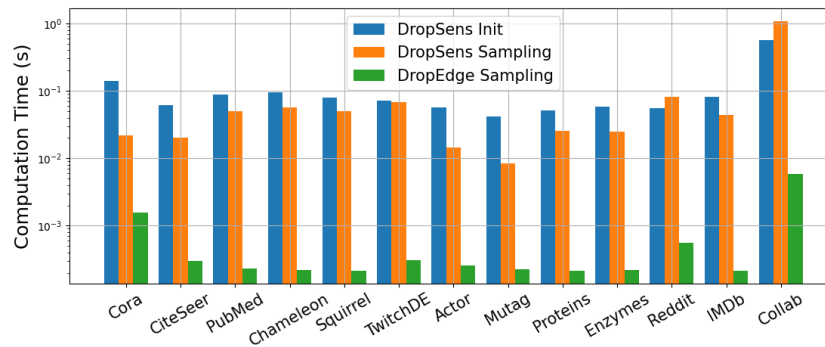


Figure 8: Initialization and sampling time of DropSens ($c = 0.8$, $q_{\max} = 0.5$) compared to sampling time of DropEdge, averaged over 10 runs. For graph classification tasks, edge masks are computed in one go (instead of one mini-batch at a time, as in practice).

Figure 7a shows the DropSens probabilities for masking incoming messages based on the in-degree of the target nodes, along with the corresponding approximations. It is clear to see that the approximation gets increasingly more accurate for higher proportions of information preserved. Figure 7b shows that the approximation is several orders of magnitude cheaper than the exact computation, and also does not scale with the in-degree d_i , as expected.⁷

Figure 8 shows the initialization time of DropSens, and compares the sampling times of DropSens and DropEdge. It is clear to see that DropSens sampling is more expensive than DropEdge, primarily because the edge-wise dropping probabilities need to be computed for new graphs,⁸ before sampling edge masks. Moreover, in our experiments, sampling consumed a very small fraction of the total training time.

F Supplementary Experiments

F.1 Test Accuracy versus DropEdge Probability

In Section 3, we studied the effect of edge-dropping probability on sensitivity between nodes at different distances. However, this analysis may be insufficient to precisely predict the impact on model performance since DropEdge-variants significantly affect the optimization trajectory as well.

⁷Note that this comparison is only valid for our SymPy-based implementation, and can be significantly different for other implementations.

⁸This step can be sped up for node-level tasks by caching edge-wise dropping probabilities.

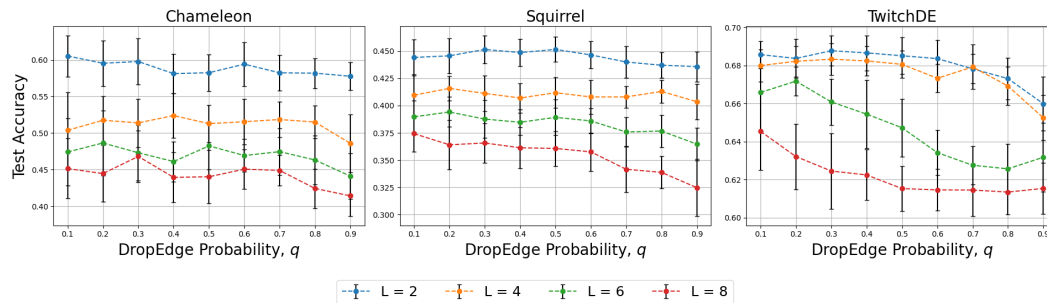


Figure 9: Dropping probability versus test accuracy of DropEdge-GCN. The theory the explains the contrasting trends as follows: random edge-dropping pushes models to fit to local information during training, which is suitable for short-range tasks, but harms test-time performance in long-range ones.

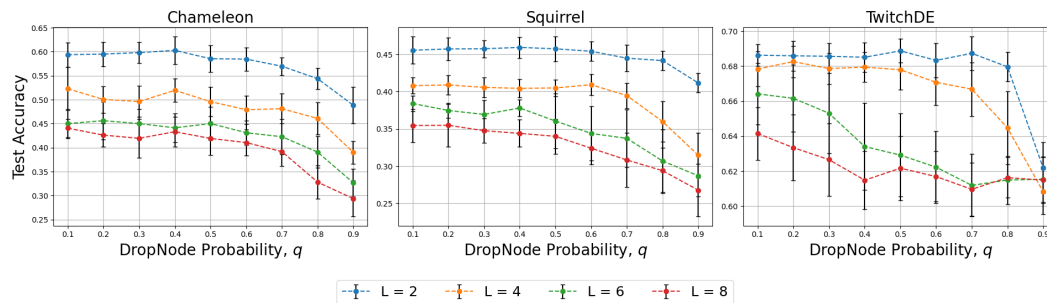


Figure 10: Dropping probability versus test accuracy of DropNode-GCN.

To learn more about the relationship between test-time performance and dropping probability, we evaluate DropEdge-GCNs on the heterophilic datasets; the results are shown in Figure 9. Clearly, on Chameleon, Squirrel and TwitchDE, the performance degrades with increasing dropping probability, as was suggested by Theorem 3.1 and Figure 1a. Surprisingly, the trends are significantly monotonic with GCNs of all depths, $L = 2, 4, 6, 8$.

F.2 Remark on DropNode

In Equation B.15, we noted that DropNode does not suffer from loss in sensitivity. However, those results were in expectation. Moreover, our analysis did not account for the effects on the learning trajectory. In practice, a high DropNode probability would make it hard for information in the node features to reach distant nodes. This would prevent the model from learning to effectively combine information from large neighborhoods, harming generalization. In Figure 10, we visualize the relationship between test-time performance and DropNode probability. The performance monotonically decreases with increasing dropping probability, as was observed with DropEdge.

F.3 Over-squashing or Under-fitting?

The results in the previous subsection suggest that using random edge-dropping to regularize model training leads to poor test-time performance. We hypothesize that this occurs because the models struggle to propagate information over long distances, causing node representations to overfit to local neighborhoods. However, a confounding effect is at play: DropEdge variants reduce the generalization gap by preventing overfitting to the training set. If this regularization is too strong, the model could underfit, which could also explain the poor test-time performance on heterophilic datasets. This concern is particularly relevant because the heterophilic networks are much larger than homophilic ones (see Table 4), making them more prone to underfitting. To investigate this, we plot the training accuracies of DropEdge-GCNs on the heterophilic datasets; Figure 11 shows the results. It is clear that the models do not underfit as the dropping probability increases. In fact, somewhat unexpectedly, the training metrics improve. Together with the results in Figure 9, we con-

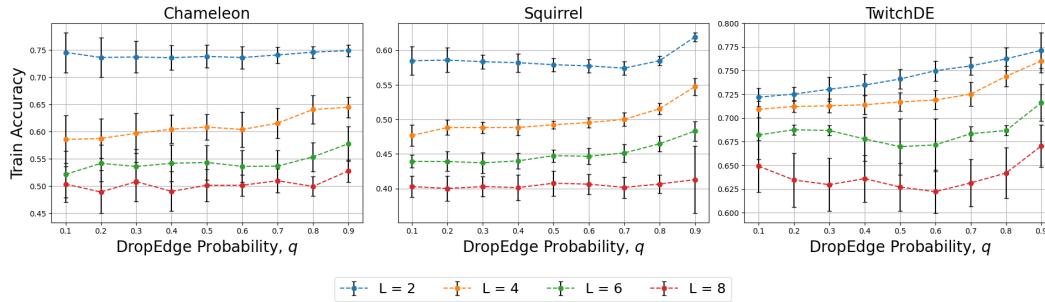


Figure 11: DropEdge probability versus training accuracy of GCNs. The training performance improves with q , suggesting that the models are not underfitting. Instead, the reason for poor test-time performance (Figure 9) is that models are over-fitting to short-range signals during training, resulting in poor generalization.

clude that DropEdge-like methods are detrimental in long-range tasks since they cause overfitting to short-range artifacts in the training data, resulting in poor generalization at test-time.

G Supplementary Experimental Results

G.1 Ranking Dropping Methods

In Table 1, we reported the mean improvement in test accuracy (%) of each dropout method over the baseline NoDrop model, using the best-performing configuration for each. Table 6 presents the corresponding raw accuracies and ranks the methods accordingly. Notably, DropSens ranks first in $7/38 \approx 18\%$ of dataset $\times$ model combinations, and places within the top 3 in $21/38 \approx 55\%$ cases, highlighting its consistent efficacy across a broad range of settings.⁹

G.2 Performance of GAT with Dropping Methods

In Table 7, we present the results of experiments in Section 5.2, but with the GAT architecture. For node classification tasks, we see the same dichotomy as in Table 1a, with dropping methods significantly improving performance on homophilic networks, while being detrimental to performance on heterophilic networks. On graph classification tasks, the dropping methods improve performance in $13/30 \approx 43\%$ cases, but the improvement, if any, is not statistically significant ($24/30 \approx 80\%$ cases). Note that the GAT architecture was unable to learn the Collab dataset, i.e. the performance in all cases was as good as a random classifier's.

G.3 Effect Size in Statistical Tests

The reliance on p-values as a measure of statistical significance has been widely criticized due to its limitations in conveying the magnitude of an effect. Although a low p-value indicates that an observed difference is unlikely to have occurred under the null hypothesis, it does not provide information about the practical significance of the result. A statistically significant effect may be too small to be meaningful in real-world applications, while a non-significant result does not necessarily imply the absence of a meaningful effect, particularly when sample sizes are small. These concerns have led to an increased emphasis on effect size measures, which quantify the magnitude of differences independently of sample size.

One widely used measure of effect size is Cohen's d statistic [14], which standardizes the difference between two group means by dividing by the pooled standard deviation:

$$d = \frac{\bar{x}_1 - \bar{x}_2}{s_p} \quad (\text{G.1})$$

$$s_p = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}} \quad (\text{G.2})$$

⁹We exclude Collab $\times$ GAT due to unsuccessful runs.

Table 6: Accuracy of dropping methods and their rankings.

GNN	Dataset	Nodrop	DropSens	DropEdge	DropNode	DropAgg	DropGNN	Dropout	DropMessage
GCN	Cora	85.06	85.17	85.27	85.27	85.15	85.43	85.48	85.20
	CiteSeer	73.39	73.81	73.90	73.09	73.94	74.11	73.13	73.48
	PubMed	83.75	83.76	83.65	83.67	83.32	83.56	84.86	84.57
	Chameleon	55.25	53.45	53.54	50.26	40.07	51.41	51.58	53.75
	Squirrel	39.74	39.72	39.40	38.43	26.00	38.29	38.44	39.19
	TwitchDE	68.01	68.26	67.97	66.83	62.71	67.06	67.50	68.28
	Actor	22.54	22.89	22.80	22.42	22.08	23.15	22.71	22.25
	Mutag	71.70	75.30	70.50	71.70	72.80	74.30	69.00	71.50
	Proteins	71.29	73.27	71.77	72.16	71.00	71.77	71.20	71.68
	Enzymes	31.39	28.42	28.36	28.31	25.97	28.49	26.91	27.81
	Reddit	82.23	80.16	75.95	71.30	64.78	69.29	81.31	78.46
	IMDb	48.58	49.34	49.26	48.26	49.12	47.40	49.16	52.00
GIN	Collab	65.27	64.16	61.92	60.76	44.33	66.15	61.97	67.88
	Cora	76.12	75.15	70.75	80.10	73.65	75.19	80.89	81.78
	CiteSeer	63.69	63.80	51.78	68.11	52.59	53.54	66.82	66.53
	PubMed	81.78	81.16	77.81	83.05	79.75	79.51	84.10	83.79
	Chameleon	44.28	34.23	36.47	41.63	37.05	34.74	43.42	48.53
	Squirrel	33.30	31.49	31.80	32.43	28.33	28.18	34.36	35.74
	TwitchDE	63.48	65.37	64.20	63.90	62.55	63.42	64.73	65.40
	Actor	20.24	20.39	22.56	20.16	23.18	20.74	20.27	20.22
	Mutag	79.50	80.90	76.60	71.00	79.70	75.80	80.10	80.60
	Proteins	69.02	68.20	67.87	69.55	69.16	68.36	69.68	70.55
	Enzymes	44.80	30.40	35.85	37.24	44.16	33.06	42.76	42.40
	Reddit	89.72	90.32	89.29	88.22	88.16	85.57	90.11	89.59
GAT	IMDb	69.96	67.50	69.72	70.34	69.16	67.22	70.04	71.30
	Collab	73.47	70.68	72.91	73.59	72.48	70.01	73.85	74.98
	Cora	83.89	84.49	84.47	84.37	84.19	84.43	84.95	84.88
	CiteSeer	72.87	73.38	73.53	72.44	73.23	73.88	72.73	72.44
	PubMed	83.36	83.45	83.28	83.00	83.02	83.05	83.44	83.54
	Chameleon	62.87	60.28	59.81	49.74	48.97	48.02	57.02	60.53
	Squirrel	44.94	41.70	41.77	36.63	36.27	35.07	38.55	42.40
	TwitchDE	67.20	67.03	66.83	64.40	66.13	65.45	66.22	67.08
	Actor	21.74	22.57	22.69	21.67	22.75	22.75	21.91	21.97
	Mutag	74.30	75.90	71.10	72.60	73.20	68.70	73.50	72.10
	Proteins	71.32	72.96	71.64	70.89	72.02	71.86	70.61	71.09
	Enzymes	28.42	29.70	28.67	25.51	26.79	25.08	26.38	29.11
	Reddit	48.71	49.31	47.96	50.15	48.87	48.35	50.23	49.90
	IMDb	48.08	49.48	47.56	49.36	48.36	47.56	49.50	50.20
	Collab	33.33	33.33	33.33	33.33	33.33	33.33	33.33	33.33

where $\bar{x}_1$ and $\bar{x}_2$ are sample means, s_1^2 and s_2^2 are unbiased sample variance, and n_1 and n_2 are the sizes of the samples. However, Cohen's d assumes that the sample standard deviation is an unbiased estimator of the population standard deviation. In small samples, this assumption does not hold, as the sample standard deviation tends to underestimate the true population variability. To address this bias, Hedges' g statistic [39] introduces a correction factor that adjusts Cohen's d statistic for small sample sizes:

$$g \approx \left(1 - \frac{3}{4(n_1 + n_2) - 9}\right) d \quad (\text{G.3})$$

[14] suggested that an effect size of 0.2 be considered small, 0.5 be considered medium, and 0.8 be considered large. In Table 8, we present Hedges' g statistic for the statistical tests in Section 5.2. We can clearly see that for homophilic datasets, there is a strong *positive effect* of using the dropping methods, but for heterophilic datasets and graph classification datasets, there is *at most* a small positive effect of using the dropping methods; rather, in most cases there is a *negative effect*.

G.4 Performance of GIN with DropSens

In Section 5.3, we showed that when modelling long-range graph-classification tasks using GCNs, DropSens outperforms state-of-the-art graph-rewiring techniques designed for alleviating over-squashing. However, it does not perform as well with GIN, as can be seen in Table 9 – unsurprisingly, since DropSens was specifically designed to work with GCN's message-passing scheme.

Table 7: Difference in mean test accuracy (%) between the best performing configuration of each dropout method and the baseline NoDrop model, with GAT as the base model. Cell colors represent p-values from a t-test evaluating whether dropout improves performance.

(a) Node classification tasks.

Dropout	Homophilic Networks			Heterophilic Networks		
	Cora	CiteSeer	PubMed	Chameleon	Squirrel	TwitchDE
DropEdge	+0.584	+0.668	−0.075	−3.059	−3.164	−0.363
DropNode	+0.476	−0.426	−0.358	−13.128	−8.303	−2.801
DropAgg	+0.303	+0.368	−0.340	−13.897	−8.669	−1.071
DropGNN	+0.543	+1.012	−0.313	−14.850	−9.869	−1.744
Dropout	+1.061	−0.140	+0.081	−5.851	−6.390	−0.981
DropMessage	+0.987	−0.425	+0.183	−2.344	−2.534	−0.114

(b) Graph classification tasks.

Dropout	Molecular Networks			Social Networks		
	Mutag	Proteins	Enzymes	Reddit	IMDb	Collab
DropEdge	−3.200	+0.321	+0.246	−0.750	−0.520	+0.000
DropNode	−1.700	−0.429	−2.914	+1.440	+1.280	+0.000
DropAgg	−1.100	+0.696	−1.632	+0.160	+0.280	+0.000
DropGNN	−5.600	+0.536	−3.340	−0.360	−0.520	+0.000
Dropout	−0.800	−0.714	−2.047	+1.520	+1.420	+0.000
DropMessage	−2.200	−0.232	+0.687	+1.190	+2.120	+0.000

This highlights the main limitation of DropSens, necessitating architecture-specific alteration to the edge-dropping strategy, which is not practical in general.

G.5 Best-performing Dropping Probabilities

For the real-world datasets in Section 5.2, we report the best performing dropping probability for different dataset–model–dropout combinations in Table 10.

Table 8: Hedges' g statistic for different dataset–model–dropout combinations. Color-coding for effect size according to [14]; red denotes negative effect, and green denotes positive effect. Medium to large positive effect sizes in bold.

		No effect		Small effect		Medium effect		Large effect	
Dataset	GNN	DropSens	DropEdge	DropNode	DropAgg	DropGNN	Dropout	DropMessage	
GCN	Cora	0.355	0.486	0.528	0.204	0.736	0.95	0.308	
	CiteSeer	0.653	0.621	−0.31	0.679	0.906	−0.29	0.121	
	PubMed	0.031	−0.298	−0.194	−1.082	−0.601	2.337	1.586	
	Chameleon	−0.782	−0.638	−1.97	−4.462	−1.465	−1.352	−0.59	
	Squirrel	−0.016	−0.222	−0.894	−8.94	−0.994	−0.952	−0.339	
	TwitchDE	0.324	−0.045	−1.275	−3.844	−0.944	−0.568	0.3	
	Actor	0.499	0.301	−0.143	−0.507	0.693	0.202	−0.322	
	Mutag	0.524	−0.126	0.0	0.121	0.292	−0.281	−0.022	
	Proteins	0.532	0.11	0.194	−0.061	0.104	−0.02	0.089	
	Enzymes	−0.57	−0.486	−0.479	−0.97	−0.462	−0.73	−0.56	
	Reddit	−0.352	−0.946	−1.826	−2.493	−2.377	−0.131	−0.588	
	IMDb	0.277	0.147	−0.081	0.121	−0.305	0.147	0.603	
	Collab	−0.314	−0.742	−1.158	−3.612	0.194	−0.775	0.658	
	Cora	−0.735	−2.983	2.424	−1.474	−0.525	3.556	3.296	
GIN	CiteSeer	0.09	−7.019	2.523	−7.432	−6.995	1.962	1.732	
	PubMed	−0.486	−1.558	1.031	−1.185	−1.447	1.961	1.71	
	Chameleon	−2.912	−1.794	−0.654	−1.635	−2.312	−0.246	0.902	
	Squirrel	−0.837	−0.499	−0.318	−1.823	−1.998	0.391	0.92	
	TwitchDE	1.056	0.35	0.17	−0.426	−0.025	0.475	0.775	
	Actor	0.258	2.094	−0.1	2.67	0.484	0.043	−0.032	
	Mutag	0.142	−0.242	−0.717	0.017	−0.347	0.049	0.097	
	Proteins	−0.196	−0.24	0.121	0.032	−0.136	0.153	0.331	
	Enzymes	−2.424	−1.272	−1.068	−0.086	−1.646	−0.293	−0.336	
	Reddit	0.313	−0.183	−0.641	−0.608	−1.563	0.17	−0.053	
	IMDb	−0.626	−0.048	0.078	−0.159	−0.527	0.015	0.252	
	Collab	−1.26	−0.203	0.048	−0.374	−1.241	0.147	0.562	
	Cora	1.233	0.992	0.829	0.496	0.894	1.842	1.788	
	CiteSeer	0.521	0.674	−0.334	0.355	1.015	−0.118	−0.319	
GAT	PubMed	0.163	−0.12	−0.63	−0.629	−0.593	0.114	0.247	
	Chameleon	−0.811	−0.782	−3.488	−3.824	−4.049	−1.334	−0.563	
	Squirrel	−1.269	−1.076	−2.964	−3.363	−3.722	−2.28	−0.799	
	TwitchDE	−0.188	−0.354	−1.741	−0.859	−0.793	−0.633	−0.093	
	Actor	1.138	1.143	−0.079	1.202	1.074	0.201	0.265	
	Mutag	0.206	−0.333	−0.184	−0.116	−0.556	−0.09	−0.234	
	Proteins	0.507	0.081	−0.109	0.162	0.136	−0.193	−0.061	
	Enzymes	0.261	0.04	−0.494	−0.26	−0.591	−0.356	0.109	
	Reddit	0.234	−0.234	0.468	0.057	−0.123	0.491	0.328	
	IMDb	0.408	−0.124	0.292	0.064	−0.134	0.318	0.474	
	Collab	0.0	0.0	0.0	0.0	0.0	0.0	0.0	

Table 9: Performance of GIN with different rewiring methods in graph-classification tasks, following [8, 44]. First, second, and third best results are colored.

Rewiring	Mutag	Proteins	Enzymes	Reddit	IMDb	Collab
None	77.70	70.80	33.80	86.79	70.18	72.99
Last FA	83.45	72.30	47.40	90.22	70.91	75.06
Every FA	72.55	70.38	28.38	50.36	49.16	32.89
DIGL	79.70	70.76	35.72	76.04	64.39	54.50
SDRF	78.40	69.81	35.82	86.44	69.72	72.96
FoSR	78.00	75.11	29.20	87.35	71.21	73.28
GTR	77.60	73.13	30.57	86.98	71.28	72.93
DropSens	80.90	68.20	30.40	90.32	67.50	70.68

Table 10: Best performing dropout configuration – $q_{\max}$ and c for DropSens, and q for other dropping methods.

GNN	Dataset	DropSens	DropEdge	DropNode	DropAgg	DropGNN	Dropout	DropMessage
GCN	Cora	0.3, 0.8	0.9	0.4	0.9	0.9	0.6	0.4
	CiteSeer	0.8, 0.9	0.7	0.2	0.8	0.9	0.4	0.1
	PubMed	0.3, 0.9	0.1	0.1	0.1	0.3	0.6	0.8
	Chameleon	0.2, 0.9	0.1	0.2	0.1	0.1	0.1	0.2
	Squirrel	0.3, 0.9	0.1	0.1	0.1	0.1	0.3	0.2
	TwitchDE	0.2, 0.95	0.2	0.2	0.1	0.1	0.1	0.2
	Actor	0.3, 0.9	0.7	0.1	0.7	0.9	0.2	0.1
	Mutag	0.2, 0.9	0.4	0.2	0.5	0.1	0.1	0.7
	Proteins	0.2, 0.95	0.2	0.3	0.5	0.5	0.2	0.3
	Enzymes	0.2, 0.95	0.1	0.1	0.1	0.1	0.2	0.3
	Reddit	0.2, 0.95	0.1	0.1	0.1	0.1	0.1	0.1
	IMDb	0.5, 0.9	0.9	0.2	0.5	0.9	0.4	0.3
	Collab	0.5, 0.95	0.2	0.1	0.1	0.1	0.1	0.1
GIN	Cora	0.2, 0.9	0.1	0.4	0.1	0.1	0.5	0.8
	CiteSeer	0.2, 0.95	0.1	0.4	0.1	0.1	0.5	0.6
	PubMed	0.2, 0.95	0.1	0.3	0.1	0.1	0.7	0.8
	Chameleon	0.2, 0.9	0.1	0.1	0.1	0.1	0.2	0.4
	Squirrel	0.2, 0.9	0.1	0.2	0.1	0.1	0.1	0.8
	TwitchDE	0.8, 0.8	0.5	0.1	0.3	0.1	0.1	0.7
	Actor	0.5, 0.8	0.9	0.4	0.9	0.9	0.7	0.2
	Mutag	0.8, 0.95	0.1	0.4	0.2	0.1	0.1	0.1
	Proteins	0.2, 0.9	0.1	0.3	0.1	0.1	0.1	0.1
	Enzymes	0.8, 0.95	0.1	0.1	0.1	0.1	0.1	0.1
	Reddit	0.2, 0.95	0.1	0.1	0.1	0.1	0.1	0.3
	IMDb	0.2, 0.9	0.1	0.3	0.3	0.1	0.1	0.6
	Collab	0.2, 0.9	0.1	0.1	0.1	0.1	0.1	0.6
GAT	Cora	0.8, 0.5	0.7	0.2	0.8	0.5	0.6	0.7
	CiteSeer	0.5, 0.95	0.9	0.1	0.9	0.9	0.1	0.1
	PubMed	0.2, 0.95	0.1	0.1	0.1	0.1	0.6	0.8
	Chameleon	0.2, 0.95	0.1	0.1	0.1	0.1	0.1	0.1
	Squirrel	0.2, 0.95	0.1	0.1	0.1	0.1	0.1	0.1
	TwitchDE	0.2, 0.95	0.1	0.1	0.1	0.6	0.1	0.1
	Actor	0.8, 0.5	0.8	0.1	0.9	0.8	0.1	0.1
	Mutag	0.5, 0.95	0.7	0.2	0.1	0.6	0.2	0.2
	Proteins	0.5, 0.8	0.7	0.1	0.8	0.5	0.1	0.4
	Enzymes	0.8, 0.95	0.1	0.1	0.1	0.3	0.2	0.2
	Reddit	0.8, 0.9	0.8	0.6	0.1	0.8	0.9	0.8
	IMDb	0.3, 0.95	0.8	0.9	0.3	0.3	0.4	0.6
	Collab	0.2, 0.9	0.1	0.1	0.1	0.1	0.1	0.1

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 following quotations are from the abstract.

- “While several Dropout-style algorithms, such as DropEdge and DropMessage, have successfully addressed over-smoothing, their impact on over-squashing remains largely unexplored.” [Appendix A.2](#) discusses the evaluation protocol adopted by several methods designed for alleviating over-smoothing and training deep GNNs, highlighting how their focus has been largely limited to homophilic networks, which are commonly short-range tasks [\[53, 100\]](#).
- In [Section 3](#), “we present theoretical results showing that DropEdge-variants reduce sensitivity between distant nodes, limiting their suitability for long-range tasks.”
- In [Section 4](#), “... we introduce DropSens, a sensitivity-aware variant of DropEdge that explicitly controls the proportion of information lost due to edge-dropping.”
- See [Sections 5.1](#) and [5.2](#) for “our experiments on long-range synthetic and real-world datasets confirm the predicted limitations of existing edge-dropping and feature-dropping methods.”
- In [Section 5.3](#), we show that “DropSens consistently outperforms graph rewiring techniques designed to mitigate over-squashing, ...”

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 limitations of our theoretical results in [Section 6](#), particularly pertaining to the class of MPNNs analysed.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.

- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

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

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

Justification: We have provided the proofs of the theoretical results for linear GCNs in [Appendix B](#), for nonlinear MPNNs in [Appendix C.1](#), and for test-time MC-averaging in [Appendix C.2](#). To the best of our knowledge, they are correct.

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: All necessary details for reproducing the plots in [Section 3](#) are provided in [Appendix D](#). The setup for the experiments in [Section 5](#) are detailed in [Appendix E](#).

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: We have attached a ZIP file containing the code, including instructions to reproduce the figures and the tables. We have also included all the necessary details for reproducing our experiments in [Appendix D](#) and [Appendix E](#). Once the review process is complete, we will include a link to the GitHub repository for this project.

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: We provide all relevant information in [Appendix E.2](#).

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 provide the statistical significance of the experimental results in Table 4 and Table 5, as well as the corresponding effect sizes in Table 8. We present the results with error bars in Figure 9. The variance in performance on the SyntheticZINC dataset is consistently low, as also observed by [33], so we avoid clutter by not including error bars in Figure 2.

Guidelines:

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: The details are provided in the introduction of Appendix E.

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: To the best of our knowledge, there are no ethical concerns regarding our submission.

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 assesses the efficacy of dropout methods on performance of GNNs in long-range tasks, and has no direct societal impacts to the best of our knowledge.

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 submission does not introduce any high risk algorithms, models, or datasets.

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: Our code is original, but relies on the PyTorch framework [67], and we only use data that is publicly available through the PyG API [24], ensuring that the license and terms of use are properly respected.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

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

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

Justification: We have attached a ZIP file containing the code, including instructions to reproduce the figures and the tables. We have also included all the necessary details for reproducing our experiments in [Appendix D](#) and [Appendix E](#). Once the review process is complete, we will include a link to the GitHub repository for this project.

Guidelines:

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

14. Crowdsourcing and research with human subjects

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

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

Justification: Our research does not involve crowdsourcing or research with human subjects.

Guidelines:

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

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

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

Answer: [NA]

Justification: Our research does not involve research with human subjects.

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 rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The project, in its entirety, was carried out without the assistance of any LLM.

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.