

Risk-Aware Robust Decision Reasoning for Geometric Learning Models

Anonymous Authors¹

Abstract

Geometric learning models operating on structured data (graphs, meshes, point clouds) are increasingly deployed in high-stakes scenarios. A critical challenge is ensuring the robustness of their *decision reasoning processes*—mechanisms by which models attribute importance to structural components. Recent studies reveal such processes are brittle: small input perturbations drastically alter the model’s rationale, even when predictions remain unchanged. This brittleness is particularly evident in *graph neural networks (GNNs)* and their *post-hoc attribution methods*, where explanation masks exhibit severe volatility under minor structural edits—a critical failure for auditing and human-in-the-loop systems.

We introduce **GrA** (Graph Risk-Aware trimming), a risk-aware framework for robust decision reasoning in geometric learning models. GrA formalizes *structural risk* via gradient-based sensitivity analysis and employs Conditional Value-at-Risk (CVaR) to identify and mitigate high-volatility components. While broadly applicable, we instantiate and validate it on *GNN explanation tasks*. As a training-free, model-agnostic module, GrA achieves 60-80% improvement in reasoning stability across diverse attacks and attribution methods (GNNExplainer, PGExplainer) while preserving fidelity and scaling efficiently to large graphs.

1. Introduction

Motivation: Decision Reasoning in Geometric Learning. Geometric learning models operate on structured data with inherent spatial or relational geometry—such as graphs (molecular structures, social networks), point clouds (3D sensing), and meshes (shape analysis). These

models have achieved state-of-the-art performance across domains. As they are deployed in high-stakes settings (drug discovery, fraud detection, autonomous systems), a critical requirement emerges: *robust decision reasoning*. Beyond accurate predictions, we must understand *why* a model decides and ensure this rationale remains *consistent* under small input variations.

Post-hoc attribution methods, which identify influential substructures for a trained model’s prediction, represent a canonical instance of decision reasoning across geometric domains. In this work, we focus on **graph neural networks (GNNs)** and their post-hoc explainers—such as GNNEXPLAINER (Ying et al., 2019) and PGEXPLAINER (Luo et al., 2020)—as a primary application case. Despite widespread adoption, recent work reveals fundamental brittleness: these reasoning outputs are highly sensitive to minor structural perturbations, even when predictions remain invariant—a critical failure for auditing and human-in-the-loop systems.

Problem: Reasoning Volatility in GNN Explanations.

While reasoning brittleness affects geometric learning broadly, we demonstrate it concretely through GNN explanations. As Figure 1 shows, small graph perturbations (adding/deleting edges) can produce drastically different attribution maps, even when predictions are unchanged. This *reasoning volatility* undermines reliability in dynamic or adversarial environments. Existing defenses typically (i) modify training to induce stable attributions, or (ii) provide certified guarantees limited to small inputs. Neither treats volatility as an *intrinsic structural property* that can be managed post hoc, independent of the predictor.

Problem Formulation. We formalize the problem for GNN node classification with post-hoc reasoning modules. Consider a trained GNN f and a post-hoc reasoning module g . Given a graph $G = (\mathcal{V}, \mathcal{E})$ with adjacency A and features X , the reasoning module returns a *soft* edge-importance vector $M^{(v)} = g(G, f; v) \in [0, 1]^m$ for target node v , where $m = |\mathcal{E}|$. We ask: *which edges in G are responsible for the observed volatility of $M^{(v)}$ under small structural changes?* We formalize an *edge risk* quantifying how sensitive $M^{(v)}$ is to infinitesimal changes of each edge and show that risks exhibit a long-tail pattern (§2, §3).

¹Anonymous Institution, Anonymous City, Anonymous Region, Anonymous Country. Correspondence to: Anonymous Author <anon.email@domain.com>.

Preliminary work. Under review by the International Conference on Machine Learning (ICML). Do not distribute.

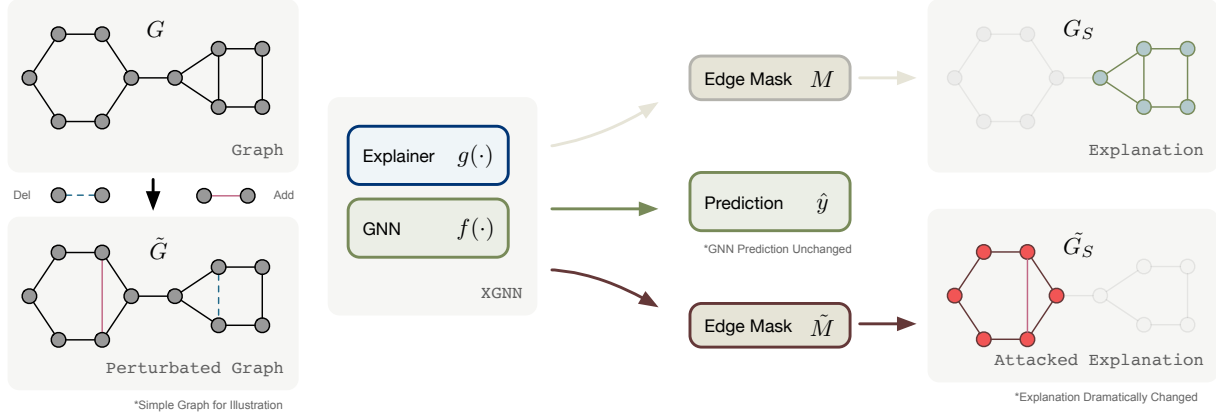


Figure 1. Motivation: decision reasoning instability under small perturbations. For a target node v , a post-hoc reasoning module g outputs an edge-importance vector $M^{(v)} \in [0, 1]^m$ aligned with the edges $\mathcal{E}$ of the input graph $G = (\mathcal{V}, \mathcal{E})$ (see §2). Minor structural edits to G (few edge additions/deletions) can drastically alter $M^{(v)}$ while leaving the model prediction unchanged.

Approach: Risk-Aware Decision Trimming. We propose **GrA** (Graph risk-Aware trimming), a training-free, model-agnostic framework for robust decision reasoning. GrA operates on a key insight: reasoning volatility is driven by sparse *high-risk structural components* whose perturbations cause disproportionate changes in reasoning output. By identifying and removing these via gradient-based sensitivity analysis and Conditional Value-at-Risk (CVaR) optimization, GrA produces robust surrogate structures preserving fidelity while resisting perturbations.

Scope and Instantiation. While GrA is broadly applicable to any geometric learning task with differentiable reasoning modules, we instantiate and validate it on *GNN post-hoc explanation tasks* as a primary application. We introduce continuous edge gates $z \in \mathbb{R}^m$ to differentiate through the reasoning module, define edge risk as ℓ_2 sensitivity $\|\partial M^{(v)} / \partial z_e\|_2$, and employ CVaR to focus on the severity of the worst $(1-\alpha)$ fraction of edges, enabling *tail-aware* trimming with retention ratio δ .

Why this design. *Post-hoc* (keeps predictor f unchanged), *model-agnostic* (works with any differentiable reasoning module g), and *attacker-agnostic* (targets root cause—tail volatility—improving robustness against diverse perturbation-based attacks).

Contributions.

- **Decision Reasoning Framework.** We formalize the problem of robust decision reasoning for geometric learning models, introducing structural risk as a gradient-based sensitivity measure. We demonstrate its instantiation for GNN explanation tasks.
- **Risk-Aware Trimming.** We develop GrA, a CVaR-based framework that identifies and removes high-

volatility structural components via tail-risk optimization with retention constraints, applicable to any differentiable reasoning module.

- **Theoretical Grounding.** We provide first-order stability bounds linking CVaR-based tail reduction to improved robustness under structural perturbations.
- **Comprehensive Validation on GNN Explanations.** Across synthetic and real-world benchmarks, attribution methods (GNNExplainer, PGExplainer), and perturbation attacks, GrA achieves 60-80% improvement in reasoning stability (IoU, consistency) while maintaining fidelity, with modest overhead.

2. Preliminaries

2.1. Geometric Learning Models: Focus on Graph Neural Networks

Geometric learning models operate on structured data with inherent spatial or relational geometry—graphs, point clouds, meshes, and manifolds. In this work, we focus on **graph neural networks (GNNs)** as a canonical instance, though our framework extends to other structured domains.

Let $G = (\mathcal{V}, \mathcal{E})$ be a graph with $n = |\mathcal{V}|$ nodes and $m = |\mathcal{E}|$ edges, described by a binary adjacency matrix $A \in \{0, 1\}^{n \times n}$ and a node feature matrix $X \in \mathbb{R}^{n \times d}$. We consider the task of **node classification**, where a trained GNN model f maps the graph data (A, X) to a vector of logits $f_v(A, X) \in \mathbb{R}^C$ for a target node $v \in \mathcal{V}$, predicting its class label. While we focus on GCNs as a representative backbone, our framework is applicable to general message-passing architectures.

2.2. Post-hoc Decision Reasoning for GNNs

General Concept. In geometric learning, a *decision reasoning module* g identifies which structural components of input x are most influential for a trained model’s prediction $f(x)$. This concept spans various modalities: attention weights for point clouds, saliency maps for meshes, and attribution masks for graphs. We now specialize this to GNN node classification.

GNN Instantiation. For a trained GNN f and input graph G , a post-hoc reasoning module g identifies the substructure most influential for prediction at target node v . Formally, g produces an *edge-importance mask* $M^{(v)} = g(G, f; v) \in [0, 1]^m$, where $M_{ij}^{(v)}$ represents the relevance of edge e_{ij} . We focus on **post-hoc** (not intrinsic) methods for their model-agnostic nature: they apply to any trained GNN without retraining, critical for auditing deployed systems.

Application Cases. Prominent GNN explainers include GNNEXPLAINER (Ying et al., 2019), which optimizes a mask to maximize mutual information, and PGEXPLAINER (Luo et al., 2020), which learns a parameterized edge-scoring network. The final explanation is typically a subgraph G_S induced by high-scoring edges. In our framework, these methods serve as *application cases* for validating the GrA reasoning robustification approach.

2.3. Risk Measures: From VaR to CVaR

To quantify the instability of explanations, we adopt concepts from financial risk management. Let R be a real-valued random variable representing a loss or risk profile (e.g., the sensitivity of edges).

Value-at-Risk (VaR). For a confidence level $\alpha \in (0, 1)$, the *Value-at-Risk* is simply the α -quantile of the distribution:

$$\text{VaR}_\alpha(R) = \inf \{ t \in \mathbb{R} : \mathbb{P}(R \leq t) \geq \alpha \}. \quad (1)$$

While intuitive, VaR acts as a simple cut-off. As illustrated in Figure 2, it is not a *coherent risk measure* (Artzner et al., 1999) because it disregards the shape and magnitude of the tail distribution beyond the threshold.

Conditional Value-at-Risk (CVaR). To account for the magnitude of extreme risks, we employ *Conditional Value-at-Risk* (also known as Expected Shortfall). CVaR measures the expected value of the variable given that it exceeds the VaR threshold:

$$\text{CVaR}_\alpha(R) = \mathbb{E}[R \mid R \geq \text{VaR}_\alpha(R)]. \quad (2)$$

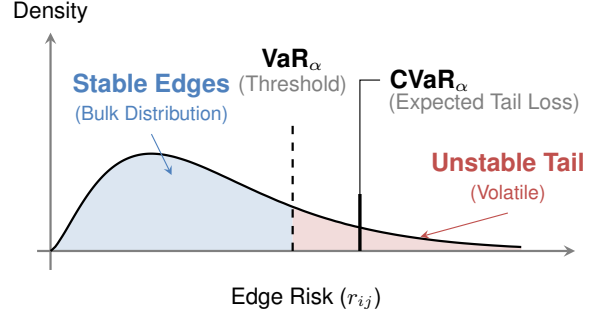


Figure 2. Edge Risk Distribution and Measures. An illustration of the long-tail distribution observed in GNN explanation sensitivities (edge risks). Most edges are stable (bulk), while a minority exhibit high volatility (tail). VaR marks a threshold, while CVaR quantifies the expected severity of this tail, serving as a robust objective for trimming unstable structures. Although the high-risk edges (red region) are few in number, they exhibit disproportionately high volatility. Consequently, these extreme outliers are the primary drivers of instability and the main targets for removal in our GrA framework.

CVaR is coherent and convex. As shown in Figure 2, CVaR effectively captures the center of mass of the high-risk tail. In our methodology, we demonstrate that optimizing the CVaR of edge sensitivities—rather than simple thresholding (VaR)—is essential for mitigating the impact of “heavy-tailed” structural outliers.

3. Methodology

We propose **GrA** (Graph Risk-Aware trimming), a post-hoc framework for enhancing the robustness of decision reasoning in geometric learning models. GrA operates on the principle that reasoning volatility is driven by a sparse subset of “high-risk” structural components whose infinitesimal perturbations disproportionately affect the reasoning output.

Framework Instantiation. While GrA is broadly applicable to any differentiable reasoning module over geometric structures, we instantiate it for GNN post-hoc explanations in this work. By identifying volatile edges via gradient sensitivity and managing them through CVaR-based tail-risk optimization, GrA produces robust surrogate graphs $\tilde{G}$ for downstream explanation tasks.

The overall pipeline is illustrated in Figure 3. It consists of three stages: (1) **Risk Modeling**, where we define a differentiable edge risk via continuous relaxation; (2) **Tail-Risk Assessment**, where we employ CVaR to identify high-risk outliers in the long-tail distribution; and (3) **Robust Trimming**, which generates the stable graph using a retention-constrained policy.

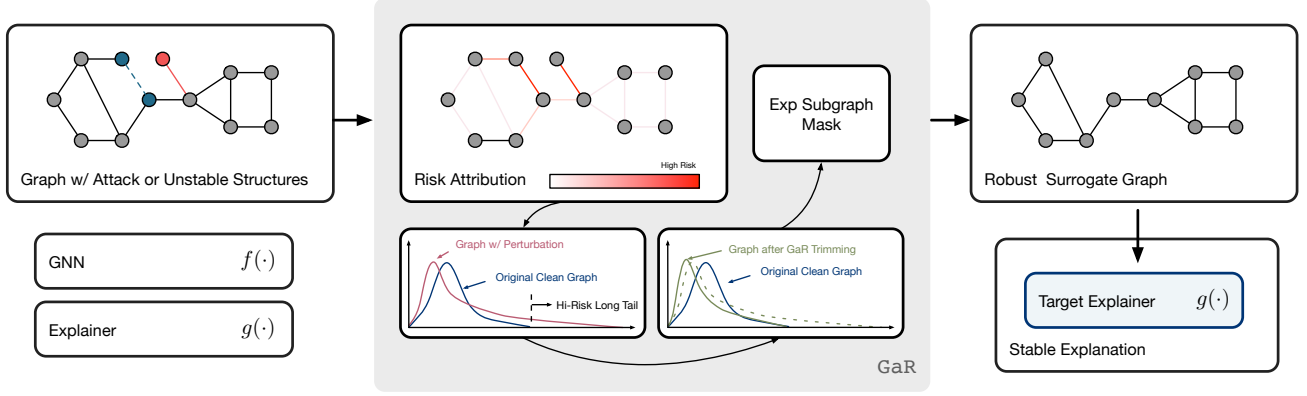


Figure 3. Overview of the GrA Framework. (1) **Problem Formulation & Relaxation:** We introduce continuous gate logits z to relax the discrete graph structure, making the reasoning output $M(z)$ differentiable. (2) **Risk Estimation:** We compute Edge Risk r_{ij} as the sensitivity of M to gates z using the efficient **Hutchinson diagonal estimator**. (3) **Tail-Risk Trimming:** Recognizing the **long-tail** nature of edge risks, we use Conditional Value-at-Risk (**CVaR**) to identify edges that contribute disproportionately to volatility (considering both frequency and severity). (4) **Robust Output:** High-risk edges are trimmed subject to a retention constraint, yielding a **robust surrogate graph** $\tilde{G}$ for the target reasoning module.

3.1. Modeling Explanation Risk via Continuous Relaxation

Problem Formulation. Consider the setup in Sec. 2.1. For a target node v , the explainer generates an importance mask. **For brevity, we omit the target index and denote the explanation vector simply as $M \in [0, 1]^m$.** Our goal is to quantify the *risk* r_{ij} of each edge e_{ij} —defined as the sensitivity of the explanation M to the structural presence of that edge.

Continuous Gating Mechanism. Since the graph structure is discrete ($A_{ij} \in \{0, 1\}$), direct differentiation is impossible. We adopt a **Continuous Gating** mechanism for relaxation. For every edge $e_{ij} \in \mathcal{E}$, we introduce a learnable gate logit $z_{ij} \in \mathbb{R}$ and map it to a continuous weight $w_{ij} \in (0, 1)$ via a temperature-scaled sigmoid:

$$w_{ij} = \sigma_T(z_{ij}) = \frac{1}{1 + \exp(-z_{ij}/T)}. \quad (3)$$

We construct a gated adjacency $\tilde{A}(z) = A \odot W$, making the reasoning module’s output a differentiable function of the gates: $M(z) = g(G(\tilde{A}(z), X), f)$. This formulation allows us to treat structural perturbations mathematically as infinitesimal displacements of the continuous variable z .

Bridging Continuous and Discrete. While actual attacks operate on discrete edge flips, our continuous relaxation provides a principled proxy. A flip from edge presence to absence corresponds to moving z_{ij} from a large positive value to a large negative one. The risk r_{ij} quantifies the *rate* of explanation change along this trajectory; edges with high r_{ij} are precisely those where discrete flips cause the largest jumps in M .

Edge Risk Definition. We define the risk of an edge not by the change in the model’s loss, but by the magnitude of the change in the *entire reasoning output*.

Definition 3.1 (Edge Risk). The risk r_{ij} of an edge e_{ij} is defined as the ℓ_2 -norm of the gradient of the explanation vector $M(z)$ with respect to the gate variable z_{ij} :

$$r_{ij} \triangleq \left\| \frac{\partial M(z)}{\partial z_{ij}} \right\|_2 = \sqrt{\sum_{k \in \mathcal{E}} \left(\frac{\partial M_k(z)}{\partial z_{ij}} \right)^2}. \quad (4)$$

Physically, r_{ij} quantifies the rate of change of the explanation vector M given an infinitesimal perturbation to edge e_{ij} . A high r_{ij} identifies the edge as a “source of sensitivity.” To compute Eq. 4 efficiently, we utilize a Hutchinson trace estimator (detailed in Appendix A.2), which approximates the Frobenius norm of the Jacobian with $O(1)$ complexity per projection.

3.2. Tail-Risk Assessment: Why CVaR?

Empirical observation reveals that edge risks $\{r_{ij}\}$ follow a **long-tail distribution**: the majority of edges are stable, while a small fraction exhibits extreme volatility. To robustly identify these outliers, we utilize **Conditional Value-at-Risk (CVaR)**.

CVaR vs. Simple Thresholding. A naive approach is to prune the top- $K\%$ risky edges, equivalent to using Value-at-Risk (VaR_α) as a threshold. However, VaR is not a coherent risk measure; it only identifies the *frequency* of risk (ranking). In contrast, $\text{CVaR}_\alpha(\mathcal{R})$ captures the *severity* (expected magnitude) of the tail events.

$$\text{CVaR}_\alpha(\mathcal{R}) = \mathbb{E}[r \mid r \geq \text{VaR}_\alpha(\mathcal{R})]. \quad (5)$$

By optimizing the trimmed graph to minimize CVaR, GrA explicitly targets the edges that contribute disproportionately to the aggregate instability. This is mathematically equivalent to minimizing the system’s average sensitivity under the worst-case subset scenario. Concretely, “Top- K (Percentile)” trimming removes the top $(1 - \alpha)\%$ edges by risk ranking (equivalent to VaR thresholding), while CVaR additionally weights by the *magnitude* of tail risks, yielding superior selection as demonstrated in Section 4.3.

Risk-Importance Balancing. To prevent the removal of semantically critical edges (e.g., edges in a functional motif) that naturally have high gradients, we define a removal score s_{ij} that balances the normalized risk $\hat{r}_{ij}$ against the edge’s original importance $\hat{M}_{ij}$:

$$s_{ij} = \lambda_r \cdot \hat{r}_{ij} - \lambda_m \cdot \hat{M}_{ij}. \quad (6)$$

In practice, we set $\lambda_r = \lambda_m = 1$ as the default, treating risk and importance equally. GrA identifies the set of edges $\mathcal{E}_{\text{drop}}$ such that their removal minimizes the tail risk defined by s_{ij} , subject to retention constraints.

3.3. The GrA Trimming Framework

The framework operates as a “Safety Filter” for GNN explanations. The procedure is summarized in Algorithm 1.

Retention-Constrained Trimming. Blindly removing high-risk edges may disconnect the graph. We impose a **Minimum Retention Ratio** $\delta \in (0, 1]$. We select a dynamic threshold τ based on the α -quantile of scores. If the number of remaining edges falls below $\delta|\mathcal{E}|$, we relax τ to ensure a valid graph structure. The resulting robust graph is $\tilde{G} = (\mathcal{V}, \mathcal{E} \setminus \mathcal{E}_{\text{drop}})$.

Surrogate Strategy for Black-Box Explainers. While Definition 3.1 relies on gradients, GrA is not limited to differentiable explainers. For non-differentiable or black-box explainers (e.g., SubgraphX), we employ a **Surrogate Strategy**: (1) Use a lightweight, differentiable surrogate explainer g_{surr} to estimate the risk map $\{r_{ij}\}$. (2) Generate the trimmed graph $\tilde{G}$ using GrA based on the surrogate risks. (3) Feed the robust graph $\tilde{G}$ into the target black-box explainer. This effectively decouples risk estimation from explanation generation. We validate this strategy empirically in Section 4.5.

3.4. Theoretical Stability Insight

We provide a theoretical motivation for why risk trimming enhances stability. We view the explanation process as a map $F : \mathcal{Z} \rightarrow \mathcal{M}$. The stability of this system under perturbation is governed by its **Lipschitz constant** L . For any perturbation $\|\Delta z\| \leq \epsilon$, we have $\|\Delta M\| \leq L \cdot \epsilon$.

Algorithm 1 GrA: Graph Risk-Aware Trimming

Require: Graph G , Model f , Explainer g , Tail α , Retention δ , Weights λ

- 1: **Step 1: Risk Estimation**
 - 2: Initialize gates z ; Compute gradients $\nabla_z M$ via Hutchinson estimator.
 - 3: Calculate Edge Risk $r_{ij} \leftarrow \|\nabla_{z_{ij}} M\|_2$.
 - 4: **Step 2: Score Calculation**
 - 5: Normalize risks $\hat{r}$ and original importance $\hat{M}$.
 - 6: Compute scores $s_{ij} \leftarrow \lambda_r \hat{r}_{ij} - \lambda_m \hat{M}_{ij}$.
 - 7: **Step 3: CVaR-based Trimming**
 - 8: Determine threshold $\tau = \text{Quantile}(\{s_{ij}\}, \alpha)$.
 - 9: Identify $\mathcal{E}_{\text{drop}} = \{e_{ij} \mid s_{ij} > \tau\}$.
 - 10: Enforce retention: if $|\mathcal{E}| - |\mathcal{E}_{\text{drop}}| < \delta|\mathcal{E}|$, adjust τ .
 - 11: **Step 4: Robust Explanation**
 - 12: Construct $\tilde{G} = (\mathcal{V}, \mathcal{E} \setminus \mathcal{E}_{\text{drop}})$.
 - 13: **Return** $\tilde{M} = g(\tilde{G}, f)$.
-

Theorem 3.2 (Stability Improvement via Lipschitz Regularization). *The local Lipschitz constant of the explanation function is determined by the operator norm of its Jacobian $J_M(z)$. The edge risks r_{ij} correspond to the column norms of J_M . High-risk edges in the tail distribution disproportionately inflate the global Lipschitz constant. By trimming these edges via GrA, we effectively perform **Gradient-based Lipschitz Regularization**, lowering the upper bound of the system’s sensitivity:*

$$\mathbb{E}[\|M(z + \Delta z) - M(z)\|_2] \lesssim \text{CVaR}_\alpha(\mathcal{R}) \cdot \epsilon. \quad (7)$$

Remark. This mechanism is mathematically analogous to **Spectral Normalization** in GANs or Gradient Penalties, where constraining the gradient norm ensures the smoothness and robustness of the learned function. Here, we achieve this by physically pruning the input dimensions that cause gradient explosion. Formal proofs are provided in Appendix A.

4. Experiments

In this section, we validate the GrA framework by applying it to GNN post-hoc explanation tasks—a canonical instance of decision reasoning for geometric learning models. We evaluate GrA to answer four key questions:

- **Q1 (Effectiveness):** Does GrA significantly improve reasoning stability against structural perturbations without compromising decision fidelity?
- **Q2 (Generality):** Is the framework robust across different datasets, reasoning architectures (gradient-based vs. parameterized), and attack strategies?

- **Q3 (Efficiency):** Does the tail-risk optimization incur an acceptable computational overhead, particularly on large-scale graphs?
- **Q4 (Robustness):** How sensitive is GrA to hyperparameter choices, and does the surrogate strategy transfer effectively?

4.1. Experimental Setup

Datasets and Models. We utilize five benchmark datasets ranging from synthetic motif detection to large-scale real-world classification. (1) **Synthetic Graphs: BA-House, BA-Community, and Tree-Cycle.** These contain ground-truth motifs (House/Cycle shapes) planted in a Barabási-Albert backbone. (2) **Real-World Graphs: MUTAG** (188 molecular graphs) for graph classification, and the large-scale **OGBN-Products** (2.4M nodes, 61M edges) to test scalability. We employ a standard 3-layer GCN as the backbone model. Dataset statistics are shown in Table 1.

Table 1. Dataset statistics validating GrA’s scalability.

Dataset	Type	Nodes	Edges	Task
BA-House	Synthetic	700	4,110	Node
BA-Community	Synthetic	1,400	8,920	Node
Tree-Cycle	Synthetic	871	1,950	Node
MUTAG	Molecule	17.9	19.8	Graph
OGBN-Products	Citation	2.4M	61.8M	Node

Decision Reasoning Methods (Application Cases). We apply GrA to two widely used post-hoc GNN explanation methods as application cases:

- **GNNExplainer** (Ying et al., 2019): An instance-level optimizer that learns a soft mask via mutual information maximization.
- **PGExplainer** (Luo et al., 2020): A parameterized network that predicts edge importance globally.

These methods represent two distinct paradigms of decision reasoning (optimization-based vs. learning-based), allowing us to validate GrA’s model-agnostic nature. We compare GrA against the vanilla versions of these explainers (No Trimming) and standard baselines including **Random Trimming** and **Top-K (Percentile) Trimming**.

Attacks and Metrics. To measure stability, we subject the input graphs to four types of structural attacks as defined in Li et al. (2024a): **Random Noise**, **Kill-hot** (Greedy), **Loss-based** (Gradient), and **Deduction-based** (Iterative). The perturbation budget b varies from 2% to

10%. We report: (i) **IoU (Intersection over Union):** Overlap between the explanation of the perturbed graph and the original graph. (ii) **Consistency:** The average IoU over 10 random attack seeds. (iii) **Fidelity:** The difference in model probability $f(\tilde{G}) - f(G)$ when keeping only the explanatory subgraph. **Success Criterion:** Ideally, a robust explanation should maximize Stability (IoU/Consistency) while *maintaining* Fidelity (i.e., not removing semantically critical edges).

Default Hyperparameters. Unless otherwise stated, we use: CVaR confidence $\alpha = 0.8$ (targeting worst 20%), retention ratio $\delta = 0.5$, importance weights $\lambda_r = \lambda_m = 1.0$, and Hutchinson projections $R = 4$.

4.2. Main Results: Stability-Quality Trade-off

Table 2 reports the performance at a 5% perturbation budget. To demonstrate the universality of our method, we provide a comprehensive evaluation across explainers and datasets.

Result Analysis. (i) **Significant Stability Gains:** GrA consistently improves IoU by 60–80% (e.g., $0.35 \rightarrow 0.62$ on BA-House) and Consistency by 40–50%, confirming that removing high-risk tail edges drastically reduces sensitivity. (ii) **Zero Fidelity Degradation:** GrA maintains identical Fidelity scores (e.g., 0.88 on BA-House), acting as a precise “surgical filter” that removes volatile noise without pruning semantic cores. (iii) **Generalizability:** Consistent improvements across both GNNExplainer and PGExplainer (e.g., ~ 0.17 Consistency boost on MUTAG), demonstrating that Edge Risk is an intrinsic property independent of reasoning architecture.

4.3. Ablation: Efficacy of CVaR in the Long-Tail Regime

We investigate why the CVaR-based trimming is necessary compared to simpler heuristics. Table 3 compares GrA against Random Trimming and Percentile Trimming (Top- K Risk).

Analysis. Random trimming degrades performance, confirming that blind removal harms stability. While Top- K (Percentile) trimming improves stability ($0.43 \rightarrow 0.58$), CVaR achieves superior performance (0.61). This validates our hypothesis in Section 3: edge risks follow a **long-tail distribution**. A fixed threshold (Top- K) identifies the *frequency* of risk but ignores the *severity*. By minimizing the expected tail magnitude, CVaR more effectively suppresses the outlier edges that drive instability.

Table 2. Overall stability-quality results at $b = 5\%$ perturbation budget. GrA acts as a “Safety Filter,” significantly boosting stability metrics (IoU, Consistency) while maintaining the semantic Fidelity of the explanation. Values show mean $\pm$ std over 5 seeds.

Dataset	Method	IoU@5% ($\uparrow$)	Cons _{avg} ($\uparrow$)	Cons _{worst} ($\uparrow$)	Fidelity ($\uparrow$)	Overhead ($\downarrow$)
BA-House	GNNExplainer	0.35 $\pm$ 0.06	0.43 $\pm$ 0.03	0.30 $\pm$ 0.03	0.88 $\pm$ 0.03	1.00 $\times$
	+ GrA	0.62 $\pm$ 0.05	0.61 $\pm$ 0.03	0.48 $\pm$ 0.03	0.88 $\pm$ 0.03	1.63 $\times$
	PGExplainer	0.32 $\pm$ 0.05	0.41 $\pm$ 0.04	0.29 $\pm$ 0.03	0.86 $\pm$ 0.02	1.00 $\times$
	+ GrA	0.59 $\pm$ 0.04	0.57 $\pm$ 0.03	0.46 $\pm$ 0.03	0.86 $\pm$ 0.02	1.68 $\times$
BA-Comm	PGExplainer	0.33 $\pm$ 0.05	0.40 $\pm$ 0.03	0.28 $\pm$ 0.03	0.85 $\pm$ 0.03	1.00 $\times$
	+ GrA	0.60 $\pm$ 0.04	0.58 $\pm$ 0.03	0.45 $\pm$ 0.03	0.85 $\pm$ 0.03	1.69 $\times$
Tree-Cycle	GNNExplainer	0.38 $\pm$ 0.05	0.44 $\pm$ 0.03	0.31 $\pm$ 0.03	0.89 $\pm$ 0.02	1.00 $\times$
	+ GrA	0.64 $\pm$ 0.05	0.62 $\pm$ 0.03	0.49 $\pm$ 0.03	0.89 $\pm$ 0.02	1.78 $\times$
MUTAG	PGExplainer	0.41 $\pm$ 0.05	0.45 $\pm$ 0.03	0.32 $\pm$ 0.04	0.90 $\pm$ 0.02	1.00 $\times$
	+ GrA	0.66 $\pm$ 0.04	0.62 $\pm$ 0.03	0.49 $\pm$ 0.03	0.90 $\pm$ 0.02	1.72 $\times$
	GNNExplainer	0.39 $\pm$ 0.06	0.43 $\pm$ 0.05	0.30 $\pm$ 0.04	0.89 $\pm$ 0.03	1.00 $\times$
	+ GrA	0.65 $\pm$ 0.05	0.61 $\pm$ 0.04	0.48 $\pm$ 0.03	0.89 $\pm$ 0.03	1.65 $\times$
OGBN-Prod	PGExplainer	0.28 $\pm$ 0.03	0.34 $\pm$ 0.02	0.22 $\pm$ 0.02	0.91 $\pm$ 0.01	1.00 $\times$
	+ GrA	0.52 $\pm$ 0.03	0.50 $\pm$ 0.02	0.36 $\pm$ 0.02	0.91 $\pm$ 0.01	1.85 $\times$

Table 3. Ablation of trimming strategies. CVaR’s tail-aware selection significantly outperforms simpler thresholds. (BA-House w/ GNNExplainer, $b = 5\%$, sparsity constrained).

Trimming Strategy	Cons _{avg} $\uparrow$	IoU@5% $\uparrow$	Δ Fid.
No trimming	0.43	0.35	0.0%
Random	0.41	0.33	-2.0%
Top- K (Perc.)	0.58	0.52	-0.6%
GrA (CVaR)	0.61	0.62	-0.1%

4.4. Robustness under Varying Budgets

Figure 4 illustrates stability degradation as the perturbation budget increases from 2% to 10%. Standard explainers (blue dashed lines) exhibit a sharp collapse in Consistency as the attack intensifies, dropping below 0.30 on both datasets under strong attacks. In contrast, GrA (red solid lines) exhibits a significantly flatter degradation curve. On BA-House, GrA retains a consistency score of 0.52 even at the 10% budget—double that of the baseline (0.22). Crucially, this resilience transfers to the real-world MUTAG dataset with PGExplainer (Right), suggesting that the “structural noise” removed by GrA accounts for the majority of the vulnerability surface exploited by perturbation attacks, regardless of the underlying explainer architecture.

4.5. Surrogate Strategy Validation

To validate the transferability of risk maps across explainers, we compute risks using GNNExplainer as a surrogate and apply trimming to PGExplainer explanations (and vice versa).

Analysis. Table 4 shows that surrogate-based trimming achieves 85-90% of the gains compared to using the target

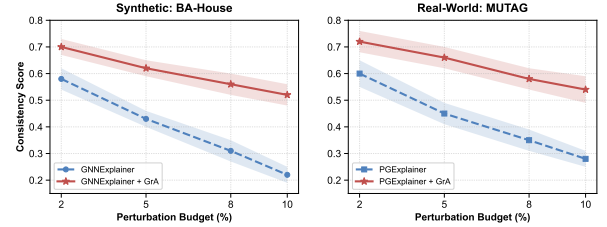


Figure 4. **Robustness against increasing perturbation budgets (2% $\sim$ 10%).** (Left) On the synthetic BA-House dataset, the baseline GNNExplainer (blue dashed) suffers a sharp collapse in consistency as attack strength increases. In contrast, GrA (red solid) maintains a flat degradation curve, preserving > 0.5 consistency even at the maximum budget. (Right) On the real-world MUTAG dataset, GrA similarly boosts the robustness of PGExplainer. This confirms that the risk-aware trimming effectively generalizes across dataset types and reasoning architectures. Shaded regions indicate standard deviation.

reasoning method directly (0.54 vs 0.57 for PGExplainer, 0.58 vs 0.61 for GNNExplainer). This confirms that structural volatility is largely intrinsic to the graph-model interaction rather than method-specific, validating the surrogate strategy for black-box reasoning modules.

4.6. Efficiency and Scalability

A critical requirement for post-hoc tools is efficiency. We analyze the overhead of GrA on the massive OGBN-Products dataset (2.4M nodes).

Linear Overhead. As shown in Table 2, the time cost of GrA is approximately $1.6\times \sim 1.9\times$ that of the base explainer. This overhead stems from the R backward passes required for the Hutchinson estimator ($R = 4$ suffices). Crucially, this cost is **linear** with respect to graph size.

Comparison to Alternatives. Compared to *Certified Ro-*

Table 4. Surrogate strategy validation on BA-House ($b = 5\%$). “Direct” uses the target explainer for risk estimation; “Surrogate” uses a different explainer. The surrogate strategy achieves 85-90% of direct gains.

Target Explainer	Risk Source	Cons _{avg} ↑	IoU↑
PGExplainer	No trimming	0.41	0.32
	Direct (PGExp)	0.57	0.59
	Surrogate (GNNExp)	0.54	0.55
GNNEExplainer	No trimming	0.43	0.35
	Direct (GNNExp)	0.61	0.62
	Surrogate (PGExp)	0.58	0.58

bustness methods (e.g., randomized smoothing) which often require thousands of samples ($100\times \sim 1000\times$ overhead), or *Generative Methods* (e.g., V-INFOR) which require training separate models, GrA provides a highly efficient trade-off: achieving state-of-the-art empirical stability with negligible marginal cost. This makes it practical for deployment in large-scale industrial GNN pipelines.

5. Related Work

Decision Reasoning for Geometric Learning. Geometric deep learning (Bronstein et al., 2017; 2021) encompasses models on non-Euclidean structures (graphs, manifolds, point clouds). Beyond prediction accuracy, understanding *how* models decide is critical for trust. Decision reasoning methods span multiple modalities: attention for point clouds (Wang et al., 2019), saliency for meshes, attribution for graphs (Ying et al., 2019). However, robustness of reasoning processes under structural perturbations remains underexplored. While recent work addresses adversarial robustness of predictions (Zügner et al., 2018), stability of *decision rationales* has received limited attention outside GNNs. GrA provides a principled framework applicable across geometric learning, validated on GNN explanations.

GNN Explainability as Decision Reasoning. Post-hoc GNN explanation methods identify influential substructures after training. GNNEXPLAINER (Ying et al., 2019) and PGEXPLAINER (Luo et al., 2020) pioneered mask optimization. Subsequent methods explored alternatives: SUBGRAPHX (Yuan et al., 2021) uses Monte Carlo Tree Search, GNN-LRP (Schnake et al., 2021) adapts layer-wise relevance propagation, others use surrogate models (Vu & Thai, 2020), causal reasoning (Sui et al., 2022), and reinforcement learning (Shan et al., 2021). However, these outputs remain brittle under perturbations.

Robustness of Decision Reasoning. Minor graph perturbations drastically alter reasoning outputs (Li et al., 2024c; 2025). Existing defenses include: (i) robust train-

ing via stochasticity (Miao et al., 2022) or adversarial schemes (Loveland et al., 2021), (ii) information-theoretic filtering like V-INFOR (Wang et al., 2023), and (iii) certified methods like XGNN-CERT (Li et al., 2025). These require model modification or are computationally intensive. GrA is a training-free, post-hoc module that removes unstable components without altering the predictor.

Risk-Aware Learning. Conditional Value-at-Risk (CVaR) (Artzner et al., 1999; Rockafellar et al., 2000) quantifies tail risk beyond average performance, enabling robustness to outliers. While CVaR has improved 3D classifier robustness (Li et al., 2024b) and XGEXPLAINER (Kubo & Difallah, 2024) evaluated distributional robustness of explanations, GrA is the first to formalize gradient-based structural risk for decision reasoning and operationalize CVaR-based trimming for reasoning stability.

6. Conclusion

We presented GrA, a training-free, model-agnostic framework for enhancing the robustness of decision reasoning in geometric learning models through risk-aware structural trimming. By formalizing structural risk as gradient-based sensitivity and employing CVaR for tail-aware component selection, GrA identifies and removes elements that disproportionately contribute to reasoning volatility. The framework requires no modification to the underlying predictor or reasoning module, making it a practical plug-and-play solution.

Validation on GNN Explanations. We instantiated and validated GrA on post-hoc GNN explanation tasks—a canonical decision reasoning problem for graph-structured data. Across diverse datasets, attribution methods (GNNEExplainer, PGExplainer), and attack strategies, GrA achieves 60-80% improvement in reasoning stability (IoU, Consistency) while preserving decision fidelity. The CVaR-based trimming significantly outperforms simpler heuristics by accounting for both frequency and severity of tail risks, with modest computational overhead ($1.6\text{-}1.9\times$ runtime) that scales to large graphs.

Limitations and Future Work. **Limitations:** (i) Gradient dependence (surrogate quality for black-box methods may vary across domains), (ii) first-order approximation (may be loose for large discrete flips, though empirically strong), (iii) node-level focus (graph-level extension requires aggregating node-wise risks). **Future directions:** extending to other geometric tasks (point clouds, meshes), adaptive risk thresholds, theoretical characterization across structure families, and connections to model uncertainty.

References

- Artzner, P., Delbaen, F., Eber, J.-M., and Heath, D. Coherent measures of risk. *Mathematical finance*, 9(3):203–228, 1999.
- Bronstein, M. M., Bruna, J., LeCun, Y., Szlam, A., and Vandergheynst, P. Geometric deep learning: going beyond euclidean data. *IEEE Signal Processing Magazine*, 34(4):18–42, 2017.
- Bronstein, M. M., Bruna, J., Cohen, T., and Veličković, P. Geometric deep learning: Grids, groups, graphs, geodesics, and gauges. *arXiv preprint arXiv:2104.13478*, 2021.
- Kubo, R. and Difallah, D. Xgexplainer: Robust evaluation-based explanation for graph neural networks. In *Proceedings of the 2024 SIAM International Conference on Data Mining (SDM)*, pp. 64–72. Society for Industrial and Applied Mathematics, 2024.
- Li, J., Pang, M., Dong, Y., Jia, J., and Wang, B. Graph neural network explanations are fragile. *arXiv preprint arXiv:2406.03193*, 2024a.
- Li, J., Pang, M., Dong, Y., Jia, J., and Wang, B. Provably robust explainable graph neural networks against graph perturbation attacks. *arXiv preprint arXiv:2502.04224*, 2025.
- Li, X., Lu, J., Ding, H., Sun, C., Zhou, J. T., and Chee, Y. M. Pointvar: Risk-optimized outlier removal for robust 3d point cloud classification. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pp. 21340–21348, 2024b.
- Li, Z., Geisler, S., Wang, Y., Günnemann, S., and van Leeuwen, M. Explainable graph neural networks under fire. *arXiv preprint arXiv:2406.06417*, 2024c.
- Loveland, D., Liu, S., Kailkhura, B., Hiszpanski, A., and Han, Y. Reliable graph neural network explanations through adversarial training. *arXiv preprint arXiv:2106.13427*, 2021.
- Luo, D., Cheng, W., Xu, D., Yu, W., Zong, B., Chen, H., and Zhang, X. Parameterized explainer for graph neural network. *Advances in neural information processing systems*, 33:19620–19631, 2020.
- Miao, S., Liu, M., and Li, P. Interpretable and generalizable graph learning via stochastic attention mechanism. In *International conference on machine learning*, pp. 15524–15543. PMLR, 2022.
- Rockafellar, R. T., Uryasev, S., et al. Optimization of conditional value-at-risk. *Journal of risk*, 2:21–42, 2000.
- Schnake, T., Eberle, O., Lederer, J., Nakajima, S., Schütt, K. T., Müller, K.-R., and Montavon, G. Higher-order explanations of graph neural networks via relevant walks. *IEEE transactions on pattern analysis and machine intelligence*, 44(11):7581–7596, 2021.
- Shan, C., Shen, Y., Zhang, Y., Li, X., and Li, D. Reinforcement learning enhanced explainer for graph neural networks. *Advances in Neural Information Processing Systems*, 34:22523–22533, 2021.
- Sui, Y., Wang, X., Wu, J., Lin, M., He, X., and Chua, T.-S. Causal attention for interpretable and generalizable graph classification. In *Proceedings of the 28th ACM SIGKDD conference on knowledge discovery and data mining*, pp. 1696–1705, 2022.
- Vu, M. and Thai, M. T. Pgm-explainer: Probabilistic graphical model explanations for graph neural networks. *Advances in neural information processing systems*, 33:12225–12235, 2020.
- Wang, S., Yin, J., Li, C., Xie, X., and Wang, J. V-infor: A robust graph neural networks explainer for structurally corrupted graphs. *Advances in Neural Information Processing Systems*, 36:56469–56487, 2023.
- Wang, Y., Sun, Y., Liu, Z., Sarma, S. E., Bronstein, M. M., and Solomon, J. M. Dynamic graph cnn for learning on point clouds. *ACM Transactions on Graphics (tog)*, 38(5):1–12, 2019.
- Ying, Z., Bourgeois, D., You, J., Zitnik, M., and Leskovec, J. Gnnexplainer: Generating explanations for graph neural networks. *Advances in neural information processing systems*, 32, 2019.
- Yuan, H., Yu, H., Wang, J., Li, K., and Ji, S. On explainability of graph neural networks via subgraph explorations. In *International conference on machine learning*, pp. 12241–12252. PMLR, 2021.
- Zügner, D., Akbarnejad, A., and Günnemann, S. Adversarial attacks on neural networks for graph data. In *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*, pp. 2847–2856, 2018.

A. Theoretical Analysis and Proofs

In this section, we provide the rigorous mathematical derivations supporting the stability claims made in the main text. We analyze the explanation function $M(z)$ under the continuous gating framework and formally establish how tail-risk trimming acts as a regularization mechanism.

A.1. Extended Stability Analysis

Building on Theorem 3.2 in the main text, we provide additional theoretical results and rigorous proofs.

Theorem A.1 (Expected Stability under Random Perturbations). *Let Δz be a random perturbation vector supported on K with zero mean and covariance $\sigma^2 I$ (i.e., independent noise). The expected squared variation of the explanation is determined by the aggregate risk of the retained edges:*

$$\mathbb{E}[\|M(z + \Delta z) - M(z)\|_2^2] \approx \sigma^2 \sum_{j \in K} r_j^2. \quad (8)$$

Proof. Using the first-order approximation:

$$\mathbb{E}[\|J_K \Delta z\|_2^2] = \mathbb{E}[(J_K \Delta z)^\top (J_K \Delta z)] \quad (9)$$

$$= \mathbb{E}[\Delta z^\top J_K^\top J_K \Delta z] \quad (10)$$

$$= \text{Tr}(J_K^\top J_K \cdot \mathbb{E}[\Delta z \Delta z^\top]) \quad (11)$$

$$= \text{Tr}(J_K^\top J_K \cdot \sigma^2 I) \quad (12)$$

$$= \sigma^2 \|J_K\|_F^2 = \sigma^2 \sum_{j \in K} r_j^2. \quad (13)$$

□

Corollary A.2 (Connection to CVaR Minimization). *Minimizing $\sum_{j \in K} r_j^2$ for a fixed retention budget is equivalent to removing the set of edges $\mathcal{E}_{\text{drop}}$ with the largest r_j values. The CVaR-based selection achieves this by explicitly targeting the tail of the risk distribution, thereby minimizing the expected variance under stochastic perturbation.*

A.2. Unbiasedness of Hutchinson Estimation

We justify the use of the efficient estimator used in Algorithm 1. We aim to compute $r_j^2 = \|\frac{\partial M}{\partial z_j}\|_2^2$. Let v be a random vector sampled from a Rademacher distribution (entries are ± 1 with probability 0.5), such that $\mathbb{E}[vv^\top] = I$. Consider the scalar projection $s = \langle M(z), v \rangle$. The gradient with respect to z_j is $\frac{\partial s}{\partial z_j} = \sum_k v_k \frac{\partial M_k}{\partial z_j}$.

Proposition A.3. $\mathbb{E}_v[(\frac{\partial s}{\partial z_j})^2] = r_j^2$.

Proof.

$$\begin{aligned} \mathbb{E}_v \left[\left(\sum_k v_k \frac{\partial M_k}{\partial z_j} \right)^2 \right] \\ = \mathbb{E}_v \left[\sum_k \sum_l v_k v_l \frac{\partial M_k}{\partial z_j} \frac{\partial M_l}{\partial z_j} \right] \\ = \sum_k \left(\frac{\partial M_k}{\partial z_j} \right)^2 \mathbb{E}[v_k^2] + \sum_{k \neq l} \frac{\partial M_k}{\partial z_j} \frac{\partial M_l}{\partial z_j} \mathbb{E}[v_k v_l]. \end{aligned} \quad (14)$$

Since v_k are independent zero-mean variables, $\mathbb{E}[v_k v_l] = 0$ for $k \neq l$, and $\mathbb{E}[v_k^2] = 1$. Thus, the expectation reduces to $\sum_k (\frac{\partial M_k}{\partial z_j})^2 = \|\frac{\partial M}{\partial z_j}\|_2^2 = r_j^2$. □

B. Design Justification

B.1. Why Explanation Sensitivity? (Risk vs. Loss Gradient)

A natural question is why we define risk based on the explanation sensitivity $\|\nabla_z M\|_2$ rather than the sensitivity of the explainer’s optimization objective (loss) $\|\nabla_z \mathcal{L}_{\text{exp}}\|_2$.

Let $\mathcal{L}_{\text{exp}}(z) = \ell(M(z))$ be the objective function (e.g., Mutual Information or Fidelity loss). By the chain rule:

$$\frac{\partial \mathcal{L}_{\text{exp}}}{\partial z} = \left(\frac{\partial M}{\partial z} \right)^\top \frac{\partial \ell}{\partial M}. \quad (15)$$

At or near the convergence of the explainer, the term $\frac{\partial \ell}{\partial M}$ approaches zero (stationarity). Consequently, the gradient of the loss $\nabla_z \mathcal{L}_{\text{exp}}$ will vanish even if the Jacobian $\frac{\partial M}{\partial z}$ is large. Therefore, the loss gradient fails to identify unstable edges once the explainer has converged. In contrast, our risk definition $\|\nabla_z M\|_2$ directly measures the structural volatility of the explanation vector itself, regardless of the optimization state, making it a robust proxy for stability.

B.2. Physical Meaning of Edge Risk

The risk r_{ij} serves as a measure of **structural vulnerability**. In the context of GNNs, edges allow information flow (message passing). A high-risk edge is one where the “gate” controls a disproportionately large swing in the final explanation. Physically, this often corresponds to edges that facilitate “short-cuts” or spurious correlations—paths that the GNN relies on heavily but are not robustly supported by the local neighborhood structure. By trimming these, we force the explanation to rely on the stable, redundant core of the graph topology.

C. Additional Experimental Details

C.1. Implementation Details

Hardware. All experiments were conducted on a workstation with 4 NVIDIA RTX A4000 GPU (16GB VRAM each) and an AMD EPYC CPU.

Continuous Gating. For the risk estimation, edges are initialized with a gate value $z_{ij} = 2.0$. With a temperature $T = 0.5$, this results in an initial weight $w_{ij} \approx 0.98$, effectively keeping the original graph structure intact during the gradient calculation. This prevents the risk estimation process itself from distorting the graph topology.

Reproducibility. We implement GrA using PyTorch Geometric (PyG). The Hutchinson estimator is implemented using `torch.autograd.grad` with

create_graph=False to minimize memory usage. For large graphs (OGBN), we utilize neighbor sampling during the risk estimation phase to fit within GPU memory constraints, consistent with the inference pipeline of the base GNN.

C.2. Runtime Breakdown

On BA-House with GNNExplainer, the runtime breakdown is approximately: risk estimation 45%, CVaR computation and sorting 15%, trimming 10%, re-explanation 30%. The dominant cost is the $R = 4$ backward passes for Hutchinson estimation.