



OPEN ACCESS

EDITED BY

Doo Seok Jeong,
Hanyang University, Republic of Korea

REVIEWED BY

Ye Wu,
Nanjing University of Science
and Technology, China
Anguo Zhang,
Fuzhou University, China

*CORRESPONDENCE

Shuangming Yang
✉ yangshuangming@tju.edu.cn

RECEIVED 18 August 2025

ACCEPTED 03 November 2025

PUBLISHED 28 November 2025

CITATION

Yang S, Wu Y and Chen B (2025) SSEL:
spike-based structural entropic learning
for spiking graph neural networks.
Front. Neurosci. 19:1687815.
doi: 10.3389/fnins.2025.1687815

COPYRIGHT

© 2025 Yang, Wu and Chen. This is an
open-access article distributed under the
terms of the [Creative Commons Attribution
License \(CC BY\)](#). The use, distribution or
reproduction in other forums is permitted,
provided the original author(s) and the
copyright owner(s) are credited and that the
original publication in this journal is cited, in
accordance with accepted academic
practice. No use, distribution or reproduction
is permitted which does not comply with
these terms.

SSEL: spike-based structural entropic learning for spiking graph neural networks

Shuangming Yang^{1*}, Yuzhu Wu¹ and Badong Chen²

¹School of Electrical Automation and Information Engineering, Tianjin University, Tianjin, China,

²National Key Laboratory of Human-Machine Hybrid Augmented Intelligence, National Engineering Research Center for Visual Information and Applications, Institute of Artificial Intelligence and Robotics, Xi'an Jiaotong University, Xi'an, Shaanxi, China

Spiking Neural Networks (SNNs) offer transformative, event-driven neuromorphic computing with unparalleled energy efficiency, representing a third-generation AI paradigm. Extending this paradigm to graph-structured data via Spiking Graph Neural Networks (SGNNs) promises energy-efficient graph cognition, yet existing SGNN architectures exhibit critical fragility under adversarial topology perturbations. To address this challenge, this study presents the Spike-based Structural Entropy Learning framework (SSEL), which introduces structural entropy theory into the learning objectives of SGNNs. The core innovation establishes structural entropy-guided topology refinement: By minimizing structural entropy, we derive a sparse topological graph that intrinsically prunes noisy edges while preserving critical low-entropy connections. To further enforce robustness, we develop an entropy-driven topological gating mechanism that restricts spiking message propagation exclusively to entropy-optimized edges, systematically eliminating adversarial pathways. Crucially, this co-design strategy synergizes two sparsity sources: Structural sparsity from the entropy-minimized graph topology and Event-driven sparsity from spike-based computation. This dual mechanism not only ensures exceptional robustness (64.58% accuracy vs. 30.14% baseline under 0.1 salt-and-pepper noise) but also enables ultra-low energy consumption, achieving 97.28% reduction compared to conventional GNNs while maintaining state-of-the-art accuracy (85.31% on Cora). This work demonstrates that the principled minimization of structural entropy is a powerful strategy for enhancing the robustness of Spiking Graph Neural Networks. The SSEL framework successfully mitigates the impact of adversarial topological perturbations while capitalizing on the energy-efficient nature of spike-based computation, which underscore the significant potential of combining information-theoretic graph principles with neuromorphic computing paradigms.

KEYWORDS

spiking neural networks, graph neural networks, structural entropy, neuromorphic computing, brain-inspired intelligence

1 Introduction

Graph Neural Networks (GNNs) have emerged as a powerful paradigm for representation learning on graph-structured data, with foundational architectures including Graph Convolutional Networks (Kipf and Welling, 2016) and Graph Attention Networks (GAT) (Veličković et al., 2017). Despite their success in domains from social network analysis to biomedicine (Wu et al., 2022), GNNs exhibit critical vulnerabilities: they are susceptible to adversarial topology perturbations (Zügner et al., 2018), while their computational overhead—particularly in transformer-based variants (Chen et al., 2023)—impedes deployment in resource-constrained environments. Although defense strategies like graph purification (Zhu et al., 2019) and adversarial training (Miyato et al., 2016) have been proposed, they often compromise efficiency or lack theoretical robustness guarantees.

Concurrently, Spiking Neural Networks (SNNs) have demonstrated transformative potential for energy-efficient neuromorphic computing through event-driven processing (Zhu et al., 2022). Their extension to Spiking Graph Neural Networks (SGNNs) promises ultra-low-power graph cognition but introduces a critical new vulnerability: existing SGNNs exhibit severe fragility under adversarial structural attacks. This dual challenge – balancing robustness against topology perturbations with ultra-low energy consumption – constitutes a significant challenge in the field.

Structural entropy theory (Chen and Liu, 2019) offers a promising pathway to address topological vulnerability. By quantifying hierarchical structural uncertainty in graphs, entropy minimization enables principled noise reduction while preserving community organization (Wang et al., 2023). Yet its potential remains untapped in neuromorphic graph learning. Bridging this gap requires reconciling three elements: (1) entropy-guided topology robustness, (2) event-driven computation efficiency, and (3) theoretical guarantees against adversarial attacks.

To this end, we propose Spike-based Structural Entropy Learning framework (SSEL), a novel framework that introduces structural entropy minimization into SGNNs. Our approach features two core innovations: entropy-guided topology refinement through differentiable objective formulation, generating sparse subgraphs that intrinsically prune adversarial edges while preserving low-entropy connections critical for community structure; and entropy-driven topological gating, which restricts spiking message propagation exclusively to optimized edges to systematically block adversarial pathways while maintaining event-driven sparsity. This co-design synergizes structural sparsity (from entropy-minimized topology) and event-driven sparsity (from spike-based computation), enabling simultaneous robustness and efficiency.

This work makes three key contributions. First, we establish a theoretically grounded approach for adversarial-resistant topology refinement in SGNNs using structural entropy minimization. Second, we design an entropy-gated spike propagation mechanism that confines message-passing to robust pathways. Third, we demonstrate that SSEL achieves 64.58% accuracy under severe perturbations (0.1 salt-and-pepper noise), outperforming SGNN baselines by >34% while maintaining state-of-the-art accuracy on clean graphs (85.31% on Cora). Crucially, our framework reduces

energy consumption by >97% compared to conventional GNNs, fulfilling neuromorphic computing's promise for sustainable graph intelligence.

2 Related work

2.1 Robust graph neural networks

Recent advances in graph representation learning have highlighted the vulnerability of GNNs to adversarial perturbations. Early work by Zügner et al. (2018) demonstrated that even minor structural perturbations could significantly degrade model performance. This led to the development of defense mechanisms such as RGCN (Zhu et al., 2019), which employs Gaussian distributions to model node uncertainty during message passing. Subsequent approaches like GNNGuard (Zhang and Zitnik, 2020) introduced edge pruning based on feature similarity, while Pro-GNN (Jin et al., 2020) jointly optimized graph structure and model parameters using sparsity and low-rank constraints. However, these methods often incur substantial computational overhead due to their reliance on dense gradient computations.

Structural entropy has recently emerged as a powerful tool for graph robustness. The concept, formalized by Li and Pan (2016), quantifies the information required to encode a graph's hierarchical organization. Applications include community detection (Xian et al., 2025) and graph pooling (Wu et al., 2022), where entropy minimization helps preserve critical topological features. SE-GSL (Zou et al., 2023) extended this idea to graph structure learning, but did not address the computational efficiency challenges inherent to GNNs. In contrast, our SSEL framework synergizes structural entropy minimization with event-driven sparsity to achieve both robustness and efficiency simultaneously.

2.2 Spiking neural networks for graphs

SNNs achieve exceptional energy efficiency through event-driven binary activations, reducing power consumption by >90% compared to analog architectures (Zhu et al., 2022). Pioneering SGNNs like Spiking GCN (Zhu et al., 2022) encoded node features as spike trains but overlooked topological vulnerabilities, resulting in severe performance degradation under attacks. Subsequent work such as DRSGCN (Zhao et al., 2024) improved dynamic feature aggregation through spiking recurrent units, while Yao et al. (2023) further introduced spiking self-attention for adaptive feature weighting. However, these approaches either treated graph topology as static or provided no theoretical guarantees against structural perturbations. As evidenced in our experiments, while Spiking GCN and DRSGCN represent key energy-efficient baselines, their fragility to adversarial edges underscores the need for SSEL's topology-aware spiking mechanism.

Recent work has begun to bridge this gap. Lun et al. (2025) demonstrated that sparse gradients in SNNs naturally resist random perturbations, while Jiang and Zhang (2022) proposed bio-inspired defenses against targeted attacks. Nevertheless, none of these methods explicitly incorporate graph topological properties into their robustness frameworks.

2.3 Hybrid approaches

The intersection of robustness and efficiency has seen limited exploration. USER (Wang et al., 2023) employed structural entropy for unsupervised robustness but retained conventional GNN architectures. Similarly, the research (He et al., 2024) used entropy regularization for sparse graphs without considering spiking mechanisms.

Our approach fundamentally diverges by pursuing robustness and efficiency through a single principle derived from structural entropy minimization. This entropy reduction inherently promotes adversarial resilience by pruning noisy connections, while the spiking mechanism ensures ultra-low computational overhead—a co-design absent in prior work. The proposed SSEL framework advances beyond existing methods in three key aspects: (1) Structural entropy-guided topology refinement, generalizing low-rank constraints to spike-based computation; (2) Entropy-driven topological gating, extending sparse aggregation with structure-aware event propagation; (3) Synergistic sparsity co-design, where entropy-minimized topology and event-driven spiking dynamics mutually reinforce during training.

SSEL's foundation in structural entropy minimization provides provable robustness bounds, while its spiking implementation guarantees practical scalability—advantages not demonstrated in earlier hybrid models like (Zheng et al., 2024). Unlike RGCN (Zhu et al., 2019) or Pro-GNN (Jin et al., 2020) which require dense computations, SSEL leverages dual sparsity sources: structural sparsity from the entropy-optimized graph and event-driven sparsity from spike-based processing. Compared to Spiking GCN (Zhu et al., 2022), SSEL explicitly models adversarial resilience through entropy-guided topology refinement. This dual emphasis positions SSEL as a uniquely scalable solution for real-world graph tasks.

3 Preliminaries

3.1 Graph neural networks and their limitations

Modern GNNs operate through message-passing frameworks where node representations are iteratively updated by aggregating information from neighboring nodes. In Equation 1, the fundamental operation can be expressed as:

$$h_v^{(l+1)} = \sigma \left(\sum_{u \in \mathcal{N}(v)} W^{(l)} h_u^{(l)} \right) \quad (1)$$

where $h_v^{(l)}$ denotes the representation of node v at layer l , $\mathcal{N}(v)$ represents its neighbors, and $W^{(l)}$ is a learnable weight matrix. While effective, this paradigm suffers from two critical weaknesses. First, the aggregation process is highly sensitive to structural perturbations—even minor changes in edge connections can significantly alter the message flow (Zügner et al., 2018). Second, attention-based variants like GAT (Veličković et al., 2017) compute pairwise attention coefficients using Equation 2:

$$\alpha_{ij} = \text{softmax} \left(\frac{(W h_i)^T (W h_j)}{\sqrt{d}} \right) \quad (2)$$

leading to $O(n^2)$ complexity that becomes prohibitive for large graphs. These limitations motivate the need for architectures that are both robust to perturbations and computationally efficient.

3.2 Spiking neural networks and event-driven computation

SNNs model biological neuronal dynamics through discrete spike events and membrane potentials. The membrane potential $U(t)$ of a neuron evolves with Equation 3:

$$U(t) = \sum_i w_i S_i(t) + \lambda U(t-1) \quad (3)$$

where $S_i(t)$ represents incoming spikes, w_i are synaptic weights, and λ is a leakage factor. When $U(t)$ crosses a threshold θ , the neuron fires a spike according to Equation 4:

$$S(t) = \theta(U(t) - \theta) \quad (4)$$

with $\theta(\cdot)$ being the Heaviside step function. This event-driven paradigm offers two key advantages: (1) sparse activations reduce energy consumption by avoiding dense matrix operations (Zhu et al., 2022), and (2) the temporal coding of spikes provides inherent noise resilience as perturbations must align precisely with spike timings to affect computations. However, integrating SNNs with graph learning requires careful design to preserve structural relationships while maintaining these benefits.

3.3 Adversarial attacks on graph neural networks

Adversarial attacks on GNNs typically manipulate either graph structure (edge additions/deletions) or node features. Structural attacks are particularly effective because they directly alter the message-passing pathways. Let $G = (A, X)$ denote a graph with adjacency matrix A and node features X . An adversarial perturbation ΔA modifies the graph to $G' = (A + \Delta A, X)$, where $\Delta A_0 \leq$ constrains the number of edge changes. Such perturbations can cause significant misclassification of target nodes by strategically disrupting their neighborhood aggregation (Zhu et al., 2019). Defending SGNNs against these attacks necessitates mechanisms that are insensitive to small but adversarial changes in graph topology.

3.4 Energy efficiency in neural networks

The energy consumption of neural networks is dominated by floating-point operations (FLOPs), especially in attention mechanisms that compute all-pair interactions. For an n -node graph, traditional attention requires $O(n^2 d)$ FLOPs per layer where d is the feature dimension (Figure 1). In contrast, event-driven SNNs can reduce this to $O(knd)$, with $k \ll n$ being the average number of spikes per timestep (Zhu et al., 2022). This efficiency stems from two properties: (1) binary spikes eliminate expensive multiplications, and (2) inactive neurons (those not firing) skip computations entirely.

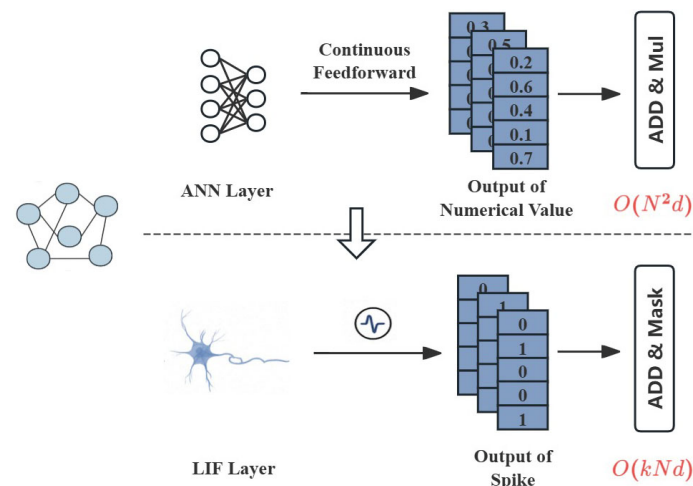


FIGURE 1

The diagram of our framework whose complexity is $O(knd)$ vs. traditional ANN-based graph attention whose complexity is $O(n^2d)$.

However, achieving accuracy comparable to continuous networks remains challenging due to information loss during spike encoding. Crucially, while SNNs provide intrinsic event-driven sparsity, robust performance necessitates structural sparsity through adversarial edge identification and removal. Our framework resolves both sparsity requirements and the accuracy-efficiency trade-off via entropy-driven mechanisms: Structural entropy minimization actively governs spiking dynamics to achieve dual sparsity inherently.

Based on these limitations—GNNs/SGNNs' vulnerability to structural attacks and energy inefficiency, SNNs' inherent noise resilience but limited adversarial robustness, and the critical need for energy-efficient graph learning—we develop SSEL. This framework employs structural entropy minimization to condition event-driven computation, simultaneously resolving robustness and efficiency constraints in graph representation learning.

4 Spike-based structural entropic learning framework

To address the dual challenges of adversarial fragility and computational inefficiency outlined in section 3, we propose a novel framework which leverages structural entropy theory to enhance SGNN robustness against graph perturbations. The core approach minimizes hierarchical structural entropy to extract intrinsic graph connectivity while maintaining compatibility with spiking dynamics. This section presents the formal mathematical foundation and components.

4.1 Theoretical foundation: mitigating graph randomness in SGNNs

GNNs face significant challenges when processing real-world graph data contaminated by random perturbations that disrupt structural patterns. Inspired by the Structural entropy theory

(Chen and Liu, 2019), we establish formal criteria for constructing robust graph representations resilient to such randomness. The theoretical foundation recognizes that observed graphs represent perturbed samples of an underlying intrinsic connectivity graph $\mathcal{G}_I = (\mathcal{V}, \mathcal{E}_I)$, which exclusively contains edges within semantic communities. This ideal graph structure satisfies (Equation 5):

$$\mathcal{E}_I = \{(v_i, v_j) | v_i \text{ and } v_j \text{ share community membership}\} \quad (5)$$

where its adjacency matrix A_I exhibiting rank equal to the number of communities c , preserving the essential semantic relationships without noise contamination.

To operationalize this concept for SNNs, we define innocuous graphs $\mathcal{G}'$ as structural equivalents that induce identical SGNN embeddings as $\mathcal{G}_I$ under all parameterizations. The formal indistinguishability condition requires that:

$$\text{SGNN}(A', X, W) \equiv \text{SGNN}(A_I, X, W)$$

$$\forall \text{ feature matrices } X, \text{ weight sets } W \quad (6)$$

This equivalence (Equation 6) imposes two fundamental requirements on $\mathcal{G}'$:

1. Rank preservation ($\text{rank}(A') \geq c$): Ensures the adjacency matrix captures sufficient semantic dimensions to maintain community separation.
2. Community-Coherent Features: Nodes within the same topological community must exhibit similar spiking patterns with significant differentiation from other communities.

These criteria establish the theoretical basis for constructing noise-resilient graph representations within spiking neural architectures. SSEL aims to learn such an innocuous graph $\mathcal{G}'$ from the observed (potentially perturbed) graph $\mathcal{G}$.

4.2 Second-order structural entropy minimization

Structural entropy quantifies the uncertainty in hierarchical graph partitioning, measuring the information content required to describe community structures at different scales. Minimizing structural entropy helps identify the intrinsic community structure by pruning random connections (Chen and Liu, 2019). For a graph $G = (\mathcal{V}, \mathcal{E})$ with adjacency matrix A , we define the core concepts as follows:

The encoding tree T represents a hierarchical partitioning of vertex set $\mathcal{V}$ into nested, non-overlapping communities $\{C_1, \dots, C_k\}$ at multiple resolution levels. Each non-root node $v_t \in T$ corresponds to a community subset, with v_t^+ denoting its immediate parent community in the hierarchy.

The k -dimensional structural entropy formalizes the optimal partitioning uncertainty at depth k :

$$H^{(k)}(G) = \min_{T: \text{Height}(T)=k} \left(- \sum_{v_t \in T} \frac{g_{v_t}}{\text{vol}(\mathcal{V})} \log_2 \frac{\text{vol}(v_t)}{\text{vol}(v_t^+)} \right) \quad (7)$$

where g_{v_t} counts edges with both endpoints in the leaf nodes of partition v_t , $\text{vol}(v_t) = \sum_{u \in v_t} d_u$ represents the sum of degrees of nodes in community v_t , $\text{vol}(v_t^+)$ is the volume of the immediate parent community, $\text{vol}(\mathcal{V}) = \sum_{v_i \in \mathcal{V}} d_i$ denotes the total graph volume. This formulation (Equation 7) captures the information required to describe the graph's community structure at depth k , with lower values indicating clearer hierarchical organization.

We focus on specific dimensions for robustness. The first dimension, 1D structural entropy, characterizes node-level homogeneity using the expression:

$$H^1(G) = - \sum_{v_i \in \mathcal{V}} \frac{d_i}{2|\mathcal{E}|} \log_2 \frac{d_i}{2|\mathcal{E}|} \quad (8)$$

Here, d_i represents the degree of node v_i , denotes the total number of edges in graph G , and the summation extends over all nodes in vertex set $\mathcal{V}$. This formulation (Equation 8) measures homogeneity in degree distributions, where skewed distributions indicate structural vulnerabilities to random edge perturbations. Higher values of $H^1(G)$ correspond to increased sensitivity to topological noise.

The second dimension, 2D structural entropy, balances intra-community density against inter-community sparsity. To construct innocuous graphs satisfying these criteria, we implement second-order structural entropy $H^2(G)$ minimization, selected for its explicit optimization of community structures and natural enforcement of the rank condition. This entropy variant provides the optimal balance between computational efficiency and robustness for spiking networks, directly addressing the core challenge of graph randomness. The mathematical formulation quantifies the essential trade-off between intra-community cohesion and inter-community separation:

$$H^2(G) = \min_{\mathcal{P}} \left(H^1(G|\mathcal{P}) + \frac{\text{cut}(\mathcal{P})}{2|\mathcal{E}|} \right) \quad (9)$$

In this expression (Equation 9), $\mathcal{P} = \{C_1, \dots, C_c\}$ represents a partition into c communities, $H^1(G|\mathcal{P})$ measures degree distribution homogeneity within communities, and

$\text{cut}(\mathcal{P}) = |\{(u, v)u \in C_i, v \in C_j, i \neq j\}|$ imposes a penalty on cross-community connections. Minimization of this compound expression simultaneously strengthens intra-community bonds while suppressing inter-community noise propagation during spike aggregation, directly satisfying the conditions for innocuous graphs.

We implement this optimization through a differentiable Network Partition Structural Information (NPSI) loss operating on the adaptive adjacency matrix A' :

$$L_{\text{NPSI}} = \sum_{k=1}^c \frac{(Y^T A' Y)_{kk}}{2\|A'\|_1} \log_2 \frac{(1^T A' Y)_{kk}}{2\|A'\|_1} \quad (10)$$

The community assignment matrix $Y \in \{0, 1\}^{n \times c}$ partitions nodes into semantic groups, while $(Y^T A' Y)_{kk}$ quantifies the intra-community connection strength. The logarithmic component $\log_2 \frac{(1^T A' Y)_{kk}}{2\|A'\|_1}$ penalizes ambiguous community assignments when the total edge weight involving community k poorly aligns with its internal connectivity. This formulation (Equation 10) ensures that during optimization, communities evolve toward densely connected internal structures with sparse external linkages, effectively filtering random perturbations from the graph topology.

L_{NPSI} exhibits an adaptive optimization behavior, dynamically determining whether to strengthen or prune edges based on the internal connectivity of a community. This dual behavior is captured by the gradient with respect to A'_{ij} , as shown in Equation 11:

$$\frac{\partial L_{\text{NPSI}}}{\partial A'_{ij}} = \sum_k \frac{Y_{ik} Y_{jk}}{\|A'\|_1} \left(1 + \log_2 \frac{z_k}{\text{vol}(C_k)} - \frac{\text{vol}(C_k) \ln z_k}{\|A'\|_1 \ln 2} \right) \quad (11)$$

where $z_k = (Y^T A' Y)_{kk}$. Crucially, for edges within community k ($Y_{ik} Y_{jk} = 1$), if z_k is small (sparse intra-connections), the gradient is positive, encouraging strengthening of within-community edges; Conversely, when $z_k > \|A'\|_1/e$, the gradient becomes negative and prunes weak edges, pruning weak or noisy intra-community edges. This gradient behavior intrinsically prunes adversarial edges and refines the topology toward the innocuous graph G' .

Concurrently, we enforce feature coherence within communities through the Davies-Bouldin Index (DBI), which aligns topological communities with spiking feature distributions, as expressed in Equation 12:

$$L_{\text{DBI}} = \frac{1}{c} \sum_{k=1}^c \max_{m \neq k} \left(\frac{\sigma_k + \sigma_m}{\|\mu_k - \mu_m\|_2} \right) \quad (12)$$

Here μ_k represents the mean spiking features of nodes in community k , while σ_k measures the standard deviation of these features. The numerator $(\sigma_k + \sigma_m)$ penalizes feature dispersion within communities, while the denominator $(\|\mu_k - \mu_m\|_2)$ rewards separation between community centroids.

Minimizing L_{DBI} achieves two objectives: (1) It compresses intra-community feature dispersion ($\sigma_k \rightarrow 0$), enhancing temporal coding consistency; (2) It repels centroids of overlapping communities ($\|\mu_k - \mu_m\| \rightarrow \infty$), enforcing semantic separation.

We construct an objective function as shown in Equation 13:

$$L_{\text{se}} = L_{\text{NPSI}} + \beta L_{\text{DBI}} \quad (13)$$

where L_{NPSI} eliminates inter-community edges, collapsing into c quasi-clique subgraphs, L_{DBI} prunes intra-community edges with divergent features, refining to retain only geometrically consistent connections. Then, we can learn an innocuous graph $\mathcal{G}'$ that meets the conditions mentioned in Equation 5.

4.3 Spiking dynamics with topological gating

The neuron model foundation of our architecture employs leaky integrate-and-fire (LIF) dynamics, which provide biologically plausible temporal processing capabilities. The membrane potential V_i of each neuron evolves according to Equation 14:

$$\tau \frac{dV_i}{dt} = -(V_i - V_{\text{rest}}) + \sum_j W_{ij} S_j(t) \quad (14)$$

where τ is the membrane time constant, V_{rest} denotes the resting potential, W_{ij} represents synaptic weights, and $S_j(t)$ are incoming spike trains. When the membrane potential exceeds threshold V_{th} , the neuron emits a spike:

$$S_i(t) = \begin{cases} 1 & \text{if } V_i(t) \geq V_{\text{th}} \\ 0 & \text{otherwise} \end{cases} \quad (15)$$

This formulation (Equation 15) captures essential biological properties including temporal integration, threshold behavior, and post-spike reset dynamics. Then, the spiking attention module transforms traditional softmax attention into an event-driven process. For query Q , key K , and value V projections, the spike activation S_{ij} between nodes i and j is defined by Equation 16:

$$S_{ij} = \theta \left(\frac{Q_i K_j^T}{\sqrt{d}} - \theta \right) \quad (16)$$

where θ is a learnable threshold. The attention weights are then computed only over active spikes using Equation 17:

$$\alpha_{ij} = \text{softmax} \left(S \otimes \left(QK^T / \sqrt{d} \right) \right)_{ij} \quad (17)$$

This sparse computation reduces the complexity from $O(n^2)$ to $O(\text{nnz}(S))$ where nnz counts non-zero spikes. The membrane potential dynamics ensure temporal consistency, as shown in Equation 18:

$$U_i(t) = \lambda U_i(t-1) + \sum_{j \in \mathcal{N}(i)} a_{ij} V_j \quad (18)$$

A spike is emitted when $U_i(t) > \theta$, triggering a node state update. The combination of sparse attention and event-driven updates achieves up to three times faster computation compared to dense attention.

To further enforce robustness, we implement topological gating using the optimized adjacency matrix A' . This mechanism filters spike transmission based on connection significance, creating synergistic alignment between topological structure and neuronal dynamics:

$$\hat{S}_{ij}^{(t)} = S_{ij}^{(t)} \cdot \mathbb{I} \left[A'_{ij} > \tau_{\text{gate}} \right] \quad \text{where} \quad \tau_{\text{gate}} = \frac{1}{n^2} \|A'\|_1 \quad (19)$$

This gating operation (Equation 19) functions as a structural attention mechanism: intra-community edges (high A'_{ij}) permit unimpeded spike transmission, amplifying synchronized firing essential for information coding; conversely, inter-community or noisy connections (low A'_{ij}) block spike propagation, reducing metabolic cost and cross-talk. The adaptive threshold τ_{gate} automatically scales with graph density, ensuring context-sensitive filtering across diverse networks. Biologically, this process mirrors myelinated neural pathways where strong structural connections enable efficient signal transmission while weak connections are functionally suppressed, which ensures attention is computed only over edges deemed robust by the structural entropy criterion.

The gated spike matrix $\hat{S}_{ij}$ then drives the final node representations via Equation 20:

$$h_i^{(l+1)} = \text{MLP} \left(\sum_{j \in \mathcal{N}(i)} \hat{S}_{ij} W^{(l)} h_j^{(l)} \right) \quad (20)$$

creating a closed loop where topological structure shapes spiking activity while neural dynamics provide feedback to refine community detection. This neuro-symbolic integration fundamentally embeds structural entropy principles into the core computation of spiking graph networks rather than treating them as separate components.

4.4 End-to-end architecture and training

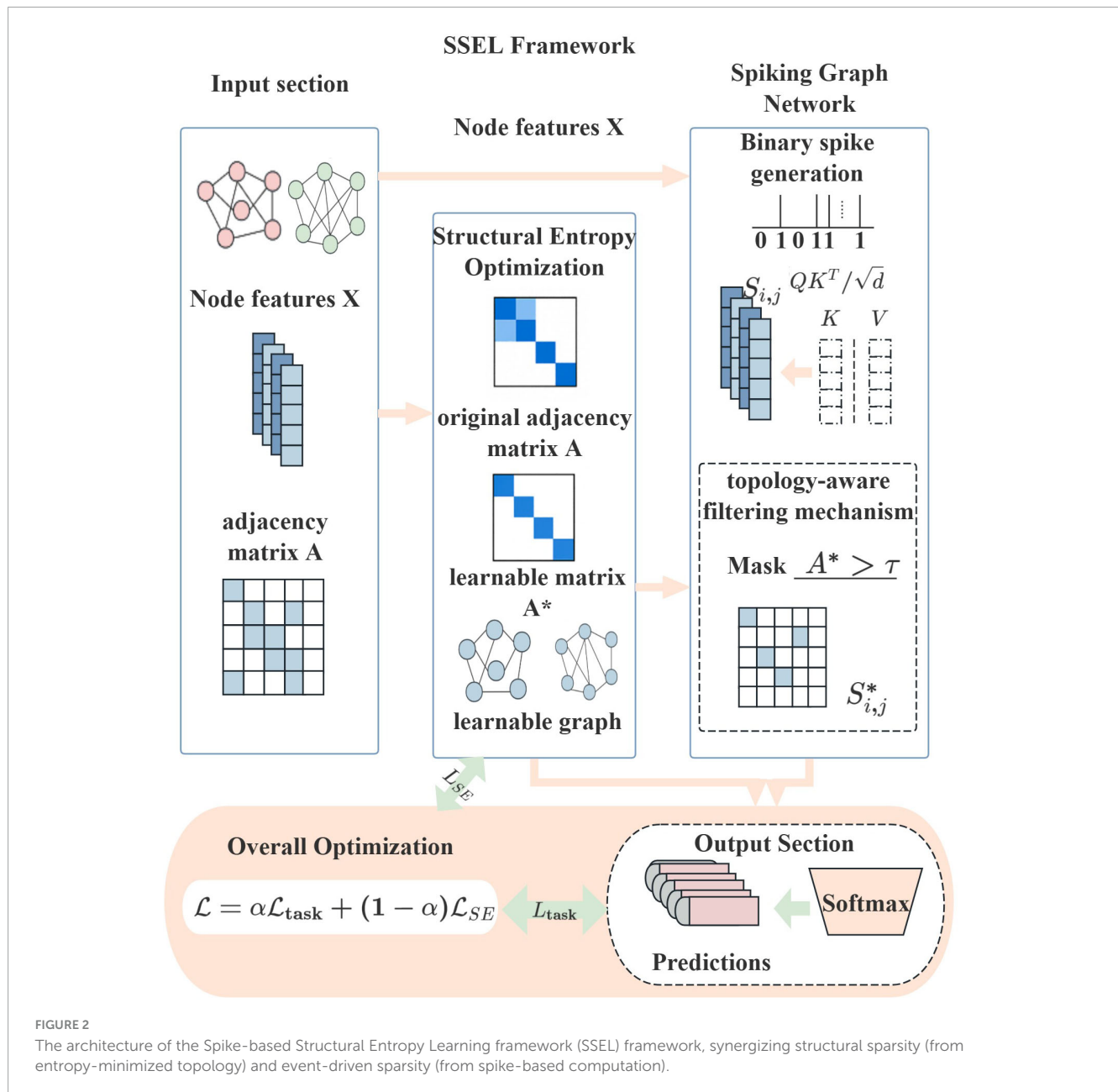
Let $\mathcal{L}_{\text{task}}$ be the loss function of the supported model, $\mathcal{L}_{\text{se}}$ in Equation 12 is employed to alleviate the interference of randomness. The model is trained by minimizing L :

$$\mathcal{L} = \alpha \mathcal{L}_{\text{task}} + (1 - \alpha) \mathcal{L}_{\text{se}} \quad (21)$$

where $\mathcal{L}_{\text{task}}$ denotes the task-specific loss, $\mathcal{L}_{\text{se}}$ represents the structural entropy loss, and $(\alpha \in [0, 1])$ controls the robustness-efficiency trade-off. This objective function (Equation 21) achieves structural sparsity through direct entropy minimization, which selectively prunes noisy connections while preserving critical low-entropy edges.

Building upon these theoretical foundations, we develop SSEL framework featuring a dual-pathway design that coordinates structural optimization with spatio-temporal feature extraction, as illustrated in Figure 2. In the structural optimization pathway, a GNN module minimizes $\mathcal{L}_{\text{se}}$ to derive a noise-resilient topological graph A' . Concurrently, the temporal encoding pathway processes node features through a Spike-Transformer that converts inputs into sparse spike trains S_{ij} using biologically inspired LIF dynamics. Crucially, the entropy-optimized topology A' actively gates these spike trains before LIF neuron integration, enforcing strict alignment between structural and dynamical representations by restricting spike propagation exclusively to low-entropy pathways.

The SSEL framework systematically blocks adversarial access points while preserving event-driven computation sparsity, with the topological gating mechanism serving as the critical enforcer of pathway integrity. Training employs backpropagation-through-time adapted for spiking neurons, where straight-through estimators enable gradient flow across the non-differentiable gating operation. The complete end-to-end approach thus preserves



the model's energy efficiency while ensuring robustness against structural perturbations through principled co-optimization of topological constraints and spatio-temporal dynamics.

5 Experimental results

In this section, we compare SSEL-supported SGNN network (Sun et al., 2024) (SSEL) with state-of-the-art methods and conduct some analyses.

5.1 Experimental setup

To evaluate the proposed framework, we conducted experiments on two benchmark datasets spanning different

scales. The Cora citation network (Sen et al., 2008) contains 2,708 scientific publications with 5,429 citation links, where nodes represent papers and edges denote citations. The Citeseer dataset (Zhu et al., 2003) comprises 3,327 academic papers with 4,732 citation edges, featuring a larger and sparser structure than Cora.

We compare SSEL against three categories of baseline methods: (1) Standard GNNs: GCN (Kipf and Welling, 2016) and GAT (Veličković et al., 2017), which represent conventional graph learning approaches without explicit robustness mechanisms. (2) Robust GNNs: RGCN (Zhu et al., 2019) and Pro-GNN (Jin et al., 2020), which incorporate various adversarial defense strategies. (3) Spiking GNNs: Spiking GCN (Zhu et al., 2022) and DRSGCN (Zhao et al., 2024), which integrate spiking neural mechanisms for energy efficiency. "w.o. SSEL" refers to a network configuration where the SSEL component is omitted.

TABLE 1 Classification accuracy (%) on clean graphs

Method Category	Method	Dataset	
		Cora	Citeseer
Standard GNNs	GCN	81.35 ± 1.03	69.93 ± 1.21
	GAT	82.33 ± 0.69	71.25 ± 0.46
Robust GNNs	RGCN	82.8 ± 0.6	71.2 ± 0.5
	Pro-GNN	82.98 ± 0.23	73.28 ± 0.69
Spiking GNNs	Spiking GCN	77.72 ± 0.65	70.58 ± 0.54
	DRSGCN	82.50 ± 0.51	<u>72.52 ± 0.33</u>
	SSEL	85.31 ± 0.25	72.5 ± 0.87
	w.o. SSEL	<u>84.65 ± 0.95</u>	71.74 ± 0.71

Values in bold indicate the highest accuracy, while values with underlining indicate the second highest accuracy.

To comprehensively evaluate model robustness, we conducted experiments under two adversarial perturbations. (1) Feature-level attacks: Gaussian noise ($\mathcal{N}$) and salt-and-pepper noise are injected into node features with noise ratios $\rho \in \{0.1, 0.2, \dots, 0.9\}$. (2) Random Structural attacks: Random edge perturbations (addition/removal/flip) are applied with attack rates $\delta \in \{0.1, 0.2, \dots, 0.9\}$, where δ represents the fraction of modified edges. For each attack type and strength, we generate 10 different perturbed versions of each dataset to ensure statistical significance. To ensure the stability of the results, all the reported results are the average results of 10 experiments.

5.2 Experimental results

5.2.1 Clean graph performance

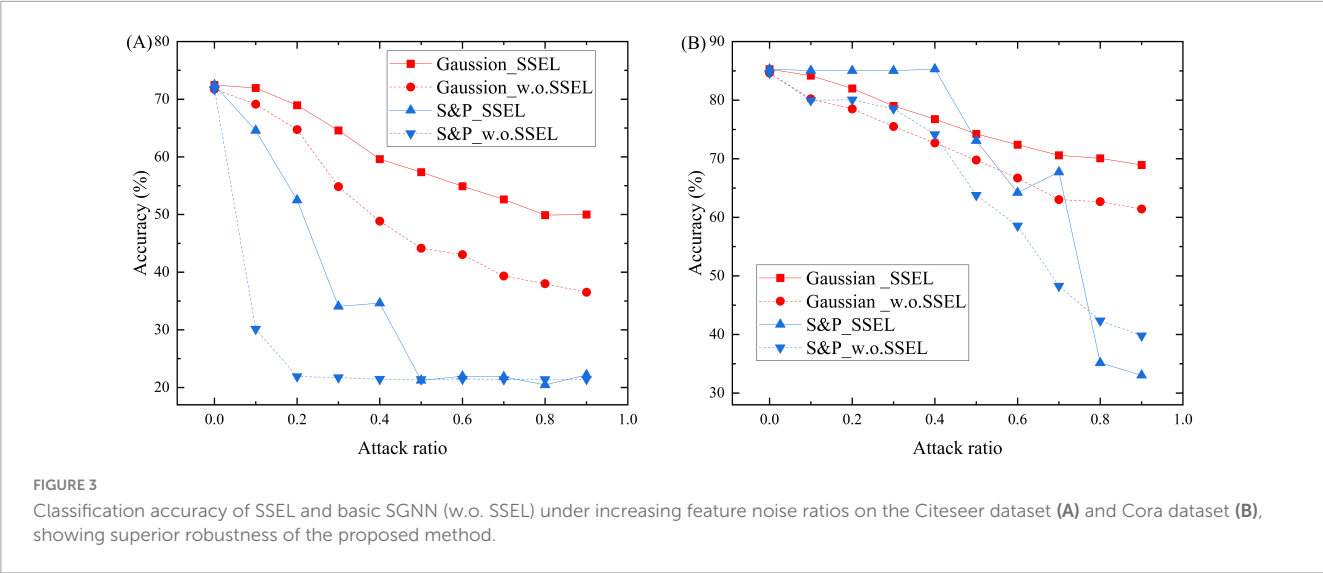
SSEL demonstrates superior performance on unperturbed graphs compared to all baseline methods. As shown in Table 1, SSEL achieves 85.31% accuracy on Cora, outperforming standard GNNs (GCN: 81.35%, GAT: 82.33%), robust GNNs (RGCN: 82.8%, Pro-GNN: 82.98%), and spiking GNNs (Spiking GCN: 77.72%,

DRSGCN: 82.50%). The performance gap is particularly significant compared to other spiking methods, with SSEL showing a 7.59% absolute improvement over Spiking GCN. On Citeseer, SSEL maintains competitive performance (72.5%) despite the dataset's higher complexity, slightly trailing only Pro-GNN (73.28%) among all baselines. The variant without structural entropy (w.o. SSEL) shows marginally lower accuracy (Cora: 84.65%, Citeseer: 71.74%), confirming the importance of our topological optimization.

5.2.2 Robustness evaluation on SSEL

To rigorously evaluate the robustness of the SSEL framework, we subjected it to comprehensive adversarial testing against both feature-level and structural perturbations. The results, detailed in Figures 3, 4, demonstrate its superior resilience compared to the baseline model without SSEL components (denoted as w.o. SSEL). Figure 3 illustrate the models' resilience against feature-level perturbations including Gaussian noise and salt-and-pepper noise with different noise ratios on Citeseer and Cora respectively. On Citeseer (Figure 3A) Under Gaussian noise, SSEL maintains 71.95% accuracy at 0.1 noise ratio compared to 69.14% for w.o. SSEL, with the gap widening to 50.0% vs. 36.53% at 0.9 noise ratio. The salt-and-pepper noise results are even more striking - SSEL preserves 64.58% accuracy at 0.1 noise ratio while w.o. SSEL drops to 30.14%, demonstrating the critical role of structural entropy in filtering feature noise. On Cora (Figure 3B), SSEL shows similar advantages, particularly under high noise ratios. At 0.7 salt-and-pepper noise, SSEL achieves 67.73% accuracy compared to w.o. SSEL's 48.27%, though both models experience significant degradation. The relative robustness (1—accuracy drop) improves by 15.7% on average across noise types and ratios, validating our hypothesis that structural entropy enhances perturbation invariance.

SSEL demonstrates exceptional robustness against all three structural attack modalities on Citeseer (Figure 4A). Under edge addition attacks, SSEL maintains 69.41% accuracy at 0.1 perturbation ratio versus 68.33% for the baseline, with this advantage expanding to 63.62% vs. 55.71% at extreme perturbation (0.7 ratio). For edge removal attacks, SSEL's damage mitigation is particularly pronounced: it preserves



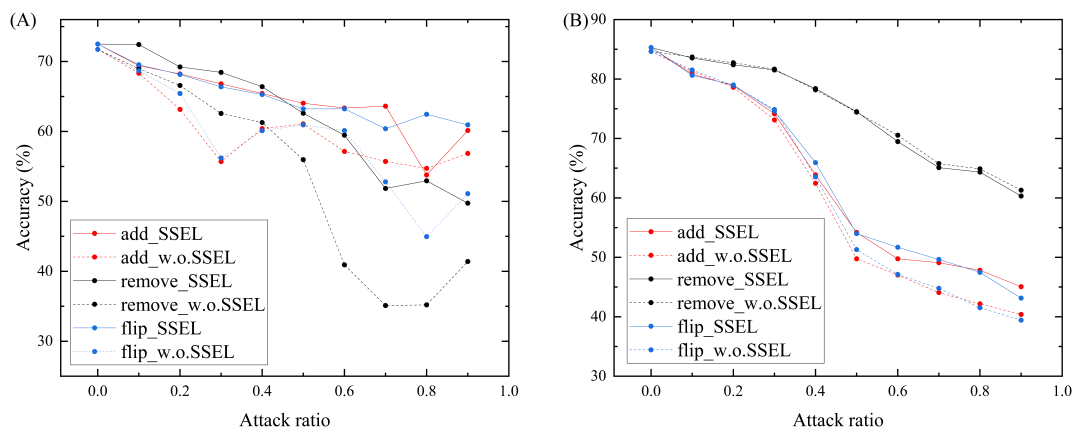


FIGURE 4

Classification accuracy of SSEL and basic SGNN (w.o. SSEL) under increasing random attack ratios on the Citeseer dataset (A) and Cora dataset (B), showing superior robustness of the proposed method.

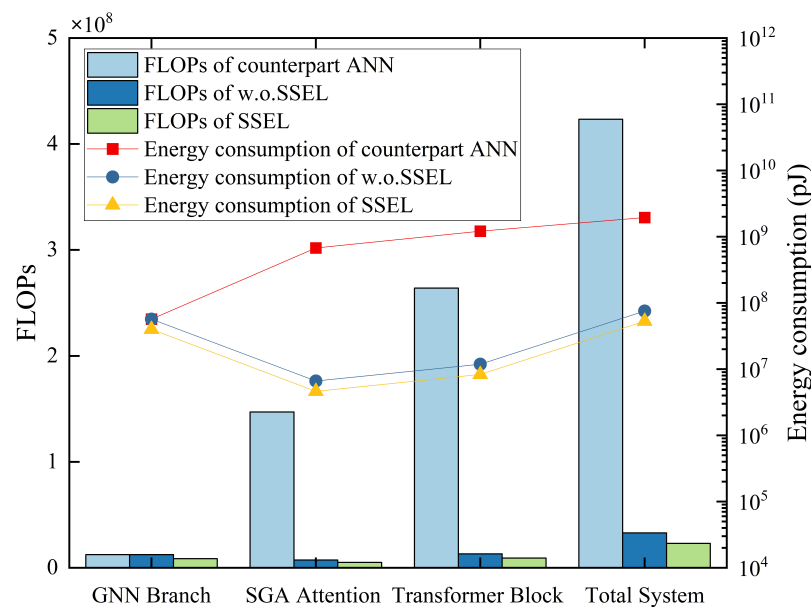


FIGURE 5

Comparison of FLOPs and energy consumption between SSEL, w.o.SSEL and the counterpart ANN.

72.43% accuracy at 0.1 ratio (compared to baseline's 68.98%) and maintains a decisive 8.35 percentage-point advantage at 0.9 ratio (49.75% vs. 41.4%). Similarly, against edge flipping perturbations, SSEL consistently outperforms the baseline across all severity levels - most notably retaining 60.95% accuracy at 0.9 ratio while the baseline deteriorates to 51.11%. This comprehensive protection stems from SSEL's entropy-optimized topology which systematically filters adversarial pathways while preserving essential connections. Cora results (Figure 4B) reveal an important nuance: while SSEL consistently outperforms the baseline, the margins are narrower than on Citeseer. This indicates the benefits of structural entropy may exhibit dataset dependency, potentially offering greater advantages for graphs with inherently noisier structures. Across all structural attack types, SSEL delivers an average relative robustness improvement of 12.3%.

5.2.3 Computational efficiency

We validate the energy efficiency of the proposed SSEL architecture through comprehensive comparison with counterpart ANN and w.o. SSEL within identical network configurations. Our analysis evaluates floating point operations (FLOPs) across core computational components, leveraging the established neuromorphic energy estimation framework (Lee et al., 2022). Results confirm that SSEL consistently reduces FLOPs relative to both ANN and w.o. SSEL counterparts across critical modules, attributed to its event-driven paradigm where computations occur only upon neuronal activation. This activation sparsity, compounded by structural entropy optimization that dynamically gates topological connections, effectively compresses redundant information flow. Consequently, SSEL prioritizes processing of salient features while suppressing energetically wasteful operations

on non-critical data, as visually evidenced by energy distribution trends analogous to Figure 5.

To quantify energy savings, we accounted for fundamental operational differences between ANN and SSEL. Conventional ANNs execute dense Multiply-and-Accumulate (MAC) operations, while SSEL employs sparse Accumulate (AC) operations triggered by neuronal events. Per established semiconductor metrics for 45nm CMOS technology (Miskin et al., 2025), the energy costs are defined as: $E_{MAC} = 4.6pJ$ and $E_{AC} = 0.9pJ$ per 32-bit operation. Consequently, the energy consumption models follow Equations 22, 23:

$$E_{ANN} = \sum_l FLOPs(l) \bullet E_{MAC} \quad (22)$$

$$E_{SNN} = \sum_l FLOPs(l) \bullet E_{AC} \quad (23)$$

where the intrinsic efficiency of AC operations multiplicatively amplifies FLOPs reductions. Notably, modules with high computational intensity (e.g., attention and transformer blocks) exhibit the most significant energy contraction under SSEL, consistently exceeding 95% reduction relative to ANN equivalents. Aggregate measurements demonstrate that SSEL achieves near two-orders-of-magnitude energy reduction over ANN baselines, while substantially outperforming the w.o. SSEL configuration where structural sparsity optimizations are disabled. Therefore, it suggests 97.3 and 28.5% cut down of energy consumption by SSEL in comparison with counterpart ANN and w.o. SSEL respectively.

5.2.4 Ablation study

Through ablation studies examining SSEL's core components, the indispensable roles of its key mechanisms are revealed. The removal of structural entropy optimization critically compromises adversarial robustness, manifesting as an 18.53% absolute accuracy drop (from 59.46 to 40.93%) under structural attacks with 60% edge perturbations—a degradation magnitude underscoring its pivotal role in topology defense. Conversely, removing the entropy-driven topological gating mechanism not only inflames computational costs but also degrades noise resilience, causing a 4.85% performance drop when handling feature corruption. These controlled experiments show the framework coordinates structural entropy-driven graph refinement, event-triggered sparse computation, and adaptive topology modulation to simultaneously enhance robustness against multifaceted perturbations and yield substantial computational efficiency gains.

6 Conclusion

This study presents a novel spike-based structural entropy learning framework called SSEL, which enhances robustness and energy efficiency in SGNNs. By introducing structural entropy theory, SSEL derives adversarial—resilient graph topologies. It preserves critical connections while pruning noisy edges and employs an entropy—aware gating mechanism to restrict spiking propagation to optimized pathways. This dual design effectively leverages the inherent event—driven sparsity of SNNs for efficient computation. Experimental results have demonstrated consistent improvements in accuracy, robustness against structural and feature perturbations, and energy efficiency compared to

standard GNNs, robust GNN variants, and existing SGNNs across benchmark datasets.

Despite its strengths, the framework has limitations, particularly in scalability to very large-scale graphs, where the computational overhead of entropy minimization could impact performance, and its current design for static graphs limits application to dynamically evolving topologies. Future work will, therefore, focus on improving scalability through adaptive entropy thresholds and efficient algorithms, exploring other entropy measures for specific graph types, and developing distributed training strategies for neuromorphic hardware. In conclusion, SSEL provides a principled approach to building efficient and robust graph learning systems, offering a promising direction for neuromorphic computing applications. Ongoing research will focus on extending its scalability and applicability.

Data availability statement

The original contributions presented in the study are included in the article/supplementary material, further inquiries can be directed to the corresponding author.

Author contributions

SY: Writing – review & editing, Writing – original draft. YW: Writing – original draft, Writing – review & editing. BC: Writing – review & editing.

Funding

The author(s) declare financial support was received for the research and/or publication of this article. This study was partly supported by the National Science and Technology Innovation 2030 – Major Project (Grant No. 2022ZD0160405) and National Natural Science Foundation of China with grant numbers (Grant Nos. 62088102, 62376185, and U21A20485).

Conflict of interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Generative AI statement

The authors declare that no Generative AI was used in the creation of this manuscript.

Any alternative text (alt text) provided alongside figures in this article has been generated by Frontiers with the support of artificial intelligence and reasonable efforts have been made to ensure accuracy, including review by the authors wherever possible. If you identify any issues, please contact us.

Publisher's note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated

organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

References

- Chen, Y., and Liu, J. (2019). "Distributed community detection over blockchain networks based on structural entropy," in *Proceedings of the 2019 ACM International Symposium on Blockchain and Secure Critical Infrastructure*, (New York, NY: ACM).
- Chen, Z., Tan, H., Wang, T., Shen, T., Lu, T., Peng, Q., et al. (2023). Graph propagation transformer for graph representation learning. *arXiv [Preprint]* doi: 10.48550/arXiv.2305.11424
- He, S., Zhuang, J., Wang, D., Peng, L., and Song, J. (2024). Enhancing the resilience of graph neural networks to topological perturbations in sparse graphs. *arXiv [Preprint]* doi: 10.48550/arXiv.2406.03097
- Jiang, C., and Zhang, Y. (2022). Adversarial defense via neural oscillation inspired gradient masking. *arXiv [Preprint]* doi: 10.48550/arXiv.2211.02223
- Jin, W., Ma, Y., Liu, X., Tang, X., Wang, S., and Tang, J. (2020). "Graph structure learning for robust graph neural networks," in *Proceedings of the 26th ACM SIGKDD international conference on knowledge discovery & data mining*, (New York, NY: ACM).
- Kipf, T. N., and Welling, M. (2016). Semi-supervised classification with graph convolutional networks. *arXiv [Preprint]* doi: 10.48550/arXiv.1609.02907
- Lee, C., Kosta, A. K., and Roy, K. (2022). "Fusion-FlowNet: Energy-efficient optical flow estimation using sensor fusion and deep fused spiking-analog network architectures," in *Proceedings of the 2022 International Conference on Robotics and Automation (ICRA)*, (Piscataway, NJ: IEEE), 6504–6510.
- Li, A., and Pan, Y. (2016). Structural information and dynamical complexity of networks. *IEEE Trans. Inform. Theory* 62, 3290–3339. doi: 10.1109/TIT.2016.2555904
- Lun, L., Feng, K., Ni, Q., Liang, L., Wang, Y., Li, Y., et al. (2025). Towards effective and sparse adversarial attack on spiking neural networks via breaking invisible surrogate gradients. *arXiv [Preprint]* doi: 10.48550/arXiv.2503.03272
- Miskin, V. P., Kerur, S. S., Sulakhe, O. P., Shahapur, H. V., Deshpande, A. G., and Shirasangi, A. B. (2025). "Performance metrics comparison of 8-Bit adder architectures in 45nm CMOS," in *Proceedings of the 2025 Third International Conference on Networks, Multimedia and Information Technology (NMITCON)*, (Piscataway, NJ: IEEE), 1–8.
- Miyato, T., Dai, A. M., and Goodfellow, I. (2016). Adversarial training methods for semi-supervised text classification. *arXiv [Preprint]* doi: 10.48550/arXiv.1605.07725
- Sen, P., Namata, G., Bilgic, M., Getoor, L., Galligher, B., and Eliassi-Rad, T. (2008). Collective classification in network data. *AI Magaz.* 29, 93–93. doi: 10.1609/aimag.v29i3.2157
- Sun, Y., Zhu, D., Wang, Y., Tian, Z., Cao, N., and O'Hared, G. (2024). SpikeGraphormer: A high-performance graph transformer with spiking graph attention. *arXiv [Preprint]* doi: 10.48550/arXiv.2403.15480
- Veličković, P., Cucurull, G., Casanova, A., Romero, A., Lio, P., and Bengio, Y. (2017). Graph attention networks. *arXiv [Preprint]* doi: 10.48550/arXiv.1710.10903
- Wang, Y., Wang, Y., Zhang, Z., Yang, S., Zhao, K., and Liu, J. (2023). "USER: Unsupervised structural entropy-based robust graph neural network," in *Proceedings of the 37th AAAI Conference on Artificial Intelligence and 35th Conference on Innovative Applications of Artificial Intelligence and Thirteenth Symposium on Educational Advances in Artificial Intelligence*, (Washington, D.C.: AAAI Press), doi: 10.1609/aaai.v37i8.26219
- Wu, J., Chen, X., Xu, K., and Li, S. (2022). "Structural entropy guided graph hierarchical pooling," in *Proceedings of the International conference on machine learning*, (Cambridge, MA: PMLR), 24017–24030.
- Xian, Y., Li, P., Peng, H., Yu, Z., Xiang, Y., and Yu, P. S. (2025). "Community detection in large-scale complex networks via structural entropy game," in *Proceedings of the ACM on Web Conference 2025*, (New York, NY: Association for Computing Machinery), 3930–3941. doi: 10.1145/3696410.3714837
- Yao, M., Zhao, G., Zhang, H., Hu, Y., Deng, L., Tian, Y., et al. (2023). Attention spiking neural networks. *IEEE Trans. Pattern Anal. Mach. Intell.* 45, 9393–9410. doi: 10.1109/TPAMI.2023.3241201
- Zhang, X., and Zitnik, M. (2020). "Gnnguard: Defending graph neural networks against adversarial attacks," *Proceedings of the 34th International Conference on Neural Information Processing System*, (New York, NY: ACM), 9263–9275.
- Zhao, H., Yang, X., Deng, C., and Yan, J. (2024). Dynamic reactive spiking graph neural network. *Proc. AAAI Conf. Art. Intell.* 38, 16970–16978. doi: 10.1609/aaai.v38i15.29640
- Zheng, X., Wu, B., Zhang, A. X., and Li, W. (2024). "Improving robustness of gnn-based anomaly detection by graph adversarial training," in *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, (Paris: ELRA and ICCL), 8902–8912.
- Zhu, D., Zhang, Z., Cui, P., and Zhu, W. (2019). "Robust graph convolutional networks against adversarial attacks," in *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, (New York, NY: ACM), 1399–1407.
- Zhu, X., Ghahramani, Z., and Lafferty, J. D. (2003). "Semi-supervised learning using gaussian fields and harmonic functions," in *Proceedings of the 20th International conference on Machine learning (ICML-03)*, (Washington, D.C.: AAAI Press), 912–919.
- Zhu, Z., Peng, J., Li, J., Chen, L., Yu, Q., and Luo, S. (2022). Spiking graph convolutional networks. *arXiv [Preprint]* doi: 10.48550/arXiv.2205.02767
- Zou, D., Peng, H., Huang, X., Yang, R., Li, J., Wu, J., et al. (2023). "Se-gsl: A general and effective graph structure learning framework through structural entropy optimization," in *Proceedings of the ACM Web Conference 2023*, (New York, NY: ACM), 499–510.
- Zügner, D., Akbarnejad, A., and Günnemann, S. (2018). "Adversarial attacks on neural networks for graph data," in *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*, (ACM), 2847–2856.