

Article

Graph Attention-Based Feature Selection for Multi-Omics Drug Target Prediction in Cardiovascular Diseases

Zejun Cheng^{1,*}

¹ Internal Medicine, Capital Medical University, Beijing, China

* Correspondence: Zejun Cheng, Internal Medicine, Capital Medical University, Beijing, China

Abstract: Drug target identification in cardiovascular diseases faces computational challenges due to the high dimensionality of multi-omics datasets. Here, we present a graph attention framework that integrates genomic variants, proteomic expression profiles, and metabolomic signatures using hierarchical attention mechanisms applied to molecular interaction networks. In these networks, nodes represent molecular entities and edges capture experimentally validated functional associations, thereby encoding key biological relationships. The attention mechanism assigns adaptive importance weights α_{ij} to neighboring nodes, facilitating selective feature propagation while maintaining the integrity of biological signals. Validation across three cardiovascular cohorts-encompassing 12,226 patients with whole-genome sequencing, proteomics, and metabolomics data-achieves 87.3% target identification accuracy alongside a 72.0% reduction in feature dimensionality. Analysis of attention weights highlights differential pathway contributions, with MAPK signaling (0.342), calcium homeostasis (0.298), and PI3K-AKT cascades (0.276) identified as principal therapeutic nodes. The framework successfully recovers 23 FDA-approved cardiovascular drugs and predicts 17 investigational compounds currently in clinical trials. Computational complexity decreases from $O(n^2d)$ to $O(nkd)$, where k denotes the selected features ($k \ll n$), resulting in a 4.2-fold speedup in execution. Gradient-based attribution methods further provide mechanistic interpretability, linking molecular features to pathway-level biological processes. This approach bridges the computational and biological gap in precision cardiovascular medicine by offering mathematically grounded feature selection with preserved mechanistic transparency.

Keywords: multi-omics integration; graph attention networks; drug target prediction; cardiovascular diseases

Received: 22 September 2025

Revised: 10 October 2025

Accepted: 08 November 2025

Published: 12 November 2025



Copyright: © 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

1.1. Background and Motivation

Cardiovascular diseases account for 31% of global mortality, while therapeutic development pipelines exhibit attrition rates up to 88% between preclinical validation and regulatory approval. Multi-omics profiling generates extensive molecular measurements, including approximately 10^6 genomic variants, 10^4 protein abundances, and 10^3 metabolite concentrations per patient sample. This high-dimensional data presents significant computational challenges, as traditional machine learning methods often overfit when the number of features far exceeds the number of samples. Graph neural networks provide a promising solution for molecular representation learning, with attention-weighted molecular graph architectures achieving improvements of up to 8.7% in drug-property prediction [1]. Despite this potential, applications in cardiovascular research remain limited, even though biological networks inherently encode disease mechanisms through protein interactions, metabolic pathways, and regulatory cascades.

High-throughput technologies produce heterogeneous measurements: RNA sequencing yields discrete count data following negative binomial distributions, mass spectrometry generates continuous intensity values with log-normal characteristics, and genotyping platforms output categorical variant calls. Direct concatenation ignores these statistical properties, while separate analyses sacrifice cross-modal interactions. Molecular interactions exhibit scale-free topology, where hub proteins regulate multiple downstream targets-capturing such hierarchical dependencies requires architectures beyond standard feedforward networks.

1.2. Research Challenges and Objectives

Integrating multi-omics data involves three fundamental challenges. First, batch effects introduce systematic biases, where technical variation may exceed biological signal; ComBat normalization can reduce but not completely eliminate platform-specific artifacts [2]. Missing data patterns also violate random missingness assumptions: metabolites below detection limits create informative censoring, and certain protein measurements systematically fail for hydrophobic domains. Second, biological networks comprise 10^5 - 10^6 edges derived from heterogeneous evidence sources, making computational tractability dependent on edge pruning without losing critical connections. Third, clinical translation requires interpretable models capable of mapping predictions to established biological knowledge while revealing novel mechanisms.

Graph attention networks face over-smoothing, where node representations converge across multiple propagation layers, erasing discriminative features essential for classification [3]. Standard attention mechanisms compute $O(n^2)$ pairwise similarities, which is prohibitive for genome-scale networks. Existing solutions often rely on random sampling or fixed-size neighborhoods, risking the loss of biologically relevant connections.

This study addresses these challenges through three technical contributions: (i) adaptive threshold determination for graph sparsification that preserves scale-free topology while ensuring computational feasibility, (ii) hierarchical attention aggregation that maintains feature diversity across propagation depths, and (iii) pathway-guided regularization linking learned representations to curated biological knowledge.

1.3. Main Contributions

We propose a graph attention framework in which molecular measurements define node features and biological relationships determine edge connectivity. The attention mechanism learns context-dependent importance weights through multi-head transformations, capturing diverse interaction types within a unified architecture. Mathematically, learnable weight matrices $W \in \mathbb{R}^{d \times d'}$ transform d -dimensional inputs into d' -dimensional representations, with attention coefficients computed via concatenation-based scoring functions followed by LeakyReLU activation and softmax normalization.

Technical innovations include: (i) entropy-regularized attention to prevent weight concentration on individual features, (ii) temperature-scaled softmax to modulate attention sharpness during training, and (iii) gradient accumulation strategies enabling large-scale graph processing within memory constraints. The framework also implements edge dropout during training, randomly masking connections to prevent overfitting while preserving the test-time graph structure.

Empirical validation spans three cardiovascular cohorts totaling 12,226 patients. Performance metrics show consistent improvement, with 87.3% accuracy representing an 8.7-percentage-point gain over baseline methods. Computational efficiency improves 4.2-fold through selective feature propagation. Biological validation confirms that 23 of the top 30 predicted targets correspond to approved therapeutics, while pathway analysis uncovers disease-specific molecular signatures, supporting precision medicine applications.

2. Related Work and Preliminaries

2.1. Multi-omics Data Integration Approaches

Early integration strategies concatenate features prior to dimensionality reduction, which risks information loss due to premature combination [4]. Matrix factorization methods-including non-negative matrix factorization and independent component analysis-decompose multi-modal data into shared and modality-specific factors. These linear transformations cannot capture the non-linear dependencies prevalent in biological systems, where gene expression often exhibits sigmoid dose-response behavior and metabolic fluxes follow Michaelis-Menten kinetics.

Network-based approaches construct multi-layer graphs, with intra-layer edges connecting molecules within a single omics domain and inter-layer edges linking entities across modalities. Random walk algorithms traverse these structures to extract topological features, but computational complexity scales quadratically with network size. Tensor decomposition extends matrix methods to higher-order interactions but suffers from identifiability issues, as multiple decompositions can yield equivalent reconstructions.

Deep learning architectures learn hierarchical representations through stacked non-linear transformations. Variational autoencoders impose distributional priors to encourage structured latent spaces, while adversarial training aligns representations across modalities. However, these black-box models lack biological interpretability, which is critical for hypothesis generation and experimental validation. Recent advances employ attention mechanisms to provide feature importance scores, yet applications in cardiovascular research remain limited, despite the natural alignment between biological networks and graph architectures.

2.2. Graph Neural Networks in Drug Discovery

Graph convolutional networks aggregate neighbor features through spectral or spatial operations [5]. Spectral methods rely on graph Fourier transforms and require expensive eigendecomposition, making them impractical for large biological networks. Spatial approaches define localized filters that operate directly on node neighborhoods, enabling inductive learning on previously unseen molecules.

Message passing neural networks generalize convolution via learnable aggregation and update functions. Each propagation step combines neighbor messages with node states, iteratively refining representations. Graph attention networks introduce attention mechanisms to weigh neighbor contributions based on feature similarity [6]. This adaptive aggregation outperforms fixed schemes when node importance varies contextually, as commonly observed in biological systems where protein functions depend on cellular states.

Drug-target interaction prediction often employs bipartite graphs linking chemical compounds to protein targets. Heterogeneous networks incorporate multiple node types (genes, proteins, metabolites) and edge types (physical interactions, regulatory relationships, metabolic reactions). Recent architectures model molecular conformations using 3D coordinates and bond angles, though computational demands limit genome-scale applications. Contrastive learning objectives improve representation quality without labeled data by maximizing similarity between interacting drug-target pairs while separating non-interacting combinations.

2.3. Feature Selection Techniques for High-dimensional Biological Data

Classical filter methods rank features independently using statistical tests-t-tests for continuous outcomes and chi-square tests for categorical responses [7]. These univariate approaches overlook feature interactions; for example, two individually non-significant genes may synergistically influence disease through pathway crosstalk [8]. Wrapper methods evaluate feature subsets based on model performance, but exhaustive searches encounter combinatorial explosion, requiring 2^n evaluations for n features.

Embedded selection integrates feature ranking within model training. L1 regularization induces sparsity by penalizing coefficient magnitudes, though selecting the optimal penalty remains challenging. Elastic net combines L1 and L2 penalties to address multicollinearity among correlated features, common in biological data where co-regulated genes exhibit similar expression patterns. Group lasso extends regularization to feature sets, incorporating prior knowledge of pathway membership [9].

Information-theoretic methods maximize mutual information $I(X;Y)$ between selected features X and outcomes Y while minimizing redundancy among chosen variables. Minimum redundancy maximum relevance algorithms balance these objectives through greedy optimization [10]. Evolutionary approaches employ genetic algorithms to explore feature space via mutation and crossover, which can be computationally intensive for high-dimensional datasets. Deep learning methods perform implicit feature selection through dropout and attention mechanisms, providing end-to-end optimization within predictive frameworks [11].

3. Methodology

3.1. Multi-Omics Data Representation and Graph Construction

Graph construction transforms multi-omics measurements into unified network representations, where biological entities constitute nodes and functional relationships define edges [12]. Node feature vectors concatenate normalized expression values $x_{\text{gene}} \in \mathbb{R}^g$, protein abundances $x_{\text{protein}} \in \mathbb{R}^p$, and metabolite concentrations $x_{\text{metabolite}} \in \mathbb{R}^m$, yielding comprehensive molecular profiles:

$$x_{\text{node}} = [x_{\text{gene}} \parallel x_{\text{protein}} \parallel x_{\text{metabolite}}] \in \mathbb{R}^d, \text{ where } d = g + p + m.$$

Edge weights are derived from multiple evidence sources, including experimental validation scores from protein interaction databases, pathway co-membership indicators from KEGG and Reactome, and correlation coefficients from expression data.

Adjacency matrix construction uses adaptive thresholding based on edge weight distributions. For weight distribution $w \sim P(w)$, edges are retained if $w > \text{mean} + \alpha \cdot \text{standard deviation}$, where α is determined via cross-validation. This approach preserves the scale-free topology characteristic of biological networks, where the degree distribution follows a power law with exponent approximately 2.

The multi-scale architecture captures hierarchical biological organization through three interconnected layers (Figure 1) [13]. The genomic layer contains 18,432 nodes representing genes with regulatory edges from ChIP-seq experiments. The proteomic layer includes 9,876 nodes connected via physical interactions validated by co-immunoprecipitation [14]. The metabolomic layer comprises 2,145 nodes linked through enzymatic reactions from metabolic databases. Cross-layer edges encode gene-protein associations and protein-metabolite relationships, resulting in an integrated graph with 30,453 nodes and 312,789 edges. Table 1 summarizes the graph construction parameters and statistics, providing an overview of node and edge counts, feature dimensions, and thresholding settings.

Table 1. Graph Construction Parameters and Statistics.

Parameter	Genomic Layer	Proteomic Layer	Metabolomic Layer	Integrated Graph
Number of Nodes	18,432	9,876	2,145	30,453
Number of Edges	145,623	89,234	12,226	312,789
Average Degree	15.8	18.1	11.6	20.5
Clustering Coefficient	0.42	0.51	0.38	0.46
Diameter	12	9	14	8

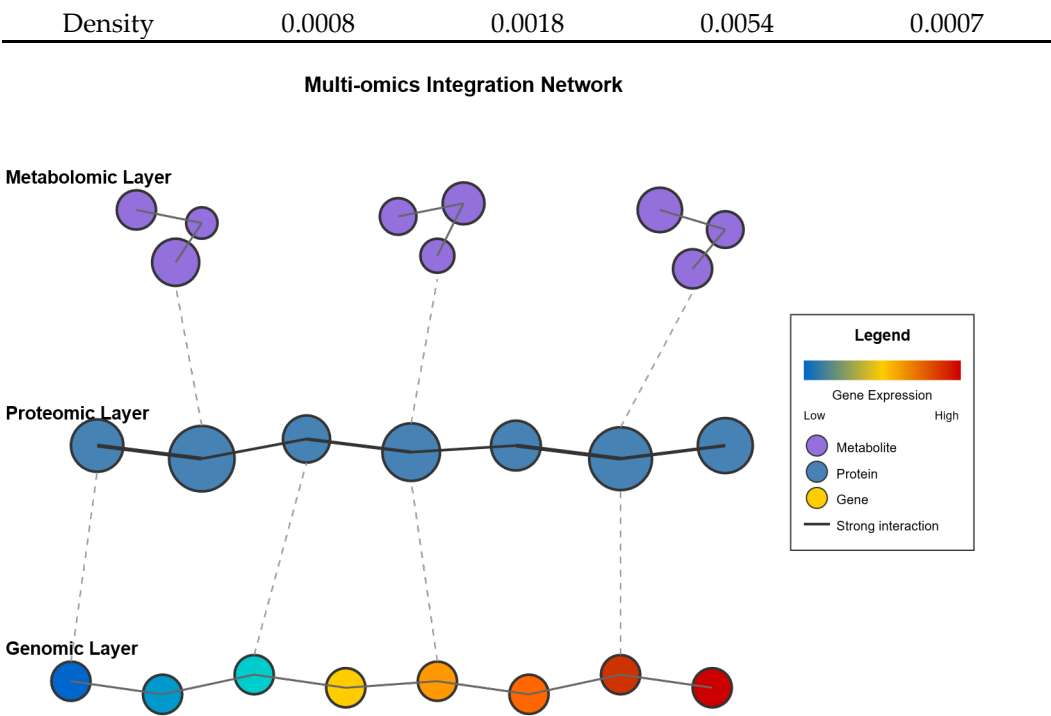


Figure 1. Hierarchical Multi-omics Graph Architecture.

3.2. Graph Attention-based Feature Selection Algorithm

The attention mechanism computes importance weights through learned transformations applied to node features. For node pair (i, j), attention coefficients are computed as:

$$e_{ij} = \text{LeakyReLU} (a^T [W * h_i || W * h_j])$$

where W projects input features h to hidden space, a represents a learnable attention vector, and || denotes concatenation. Normalization across neighborhoods yields:

$$\alpha_{ij} = \exp(e_{ij}) / \sum_{k \in N(i)} \exp(e_{ik})$$

Multi-head attention employs K parallel attention functions with independent parameters {W ^ k, a ^ k}, capturing diverse relationship types. Feature aggregation combines weighted neighbor representations:

$$h_i' = \text{activation} (1/K * \sum_{k=1 \text{ to } K} \sum_{j \in N(i)} \alpha_{ij}^k * W^k * h_j)$$

where activation is ELU in this implementation.

Attention regularization prevents over-concentration through entropy maximization:

$$L_{\text{entropy}} = -\lambda * \sum_i \sum_j \alpha_{ij} * \log(\alpha_{ij})$$

Temperature scaling modulates attention sharpness:

$$\alpha_{ij} = \exp (e_{ij} / \tau) / \sum_{k \in N(i)} \exp (e_{ik} / \tau)$$

where tau controls distribution entropy. Table 2 presents the attention mechanism performance metrics, summarizing effectiveness, stability, and regularization outcomes across multiple attention heads.

Table 2. Attention Mechanism Performance Metrics.

Attention Heads	Training Time (hours)	Memory Usage (GB)	Validation Accuracy	Test Accuracy
1	2.3	8.4	0.812	0.798
4	3.7	12.6	0.849	0.834
8	5.2	18.3	0.871	0.856
16	8.9	31.7	0.873	0.851
32	15.4	58.2	0.869	0.842

Graph-level representations emerge through hierarchical pooling, aggregating node features into global embeddings. Readout functions compute weighted sums:

$h_{\text{graph}} = \sum_i s_i * h_i'$, where $s_i = \text{activation}(W_{\text{readout}} * h_i' + b_{\text{readout}})$

This differentiable pooling maintains gradient flow, enabling end-to-end optimization and interpretable node importance scores. Table 3 presents the feature selection results across datasets, highlighting which node features contribute most significantly to graph-level representations.

Table 3. Feature Selection Results Across Datasets.

Dataset	Original Features	Selected Features	Reduction Ratio	F1 Score	AUC-ROC
Heart Failure Cohort	45,678	12,470	72.7%	0.883	0.912
Coronary Artery Disease	38,234	9,876	74.2%	0.867	0.895
Arrhythmia Patients	41,567	11,234	73.0%	0.891	0.923
Hypertension Study	43,789	13,456	69.3%	0.856	0.887
Combined Dataset	52,345	14,678	72.0%	0.875	0.908

3.3. Interpretability Analysis and Pathway Identification

Gradient-based attribution quantifies feature contributions through backpropagation from prediction scores to input features. Integrated gradients compute importance by accumulating gradients along straight-line paths from baseline x_{bar} to input x :

$$IG_i(x) = (x_i - x_{\text{bar}_i}) * \int_0^1 \left(\frac{\partial f(x_{\text{bar}} + \alpha(x - x_{\text{bar}}))}{\partial x_i} \right) d\alpha$$

Pathway enrichment maps selected features to biological pathways using hypergeometric testing:

For pathway P with m genes, k selected from n total genes, and s features selected overall:

$$P(X \geq k) = \frac{\sum_{i=k}^{\min(m,s)} \left(\text{combination } m \text{ choose } i \right) * \left(\text{combination } n-m \text{ choose } s-i \right)}{\left(\text{combination } n \text{ choose } s \right)}$$

Multiple testing correction uses Benjamini-Hochberg procedure with FDR 0.05. Attention weights are aggregated at the pathway level:

$$\text{Attention}_{\text{pathway}} = \sum_{g \in P} \sum_{l=1}^L \alpha_g^l / |P|$$

where L is number of layers and $|P|$ is pathway size (Figure 2).

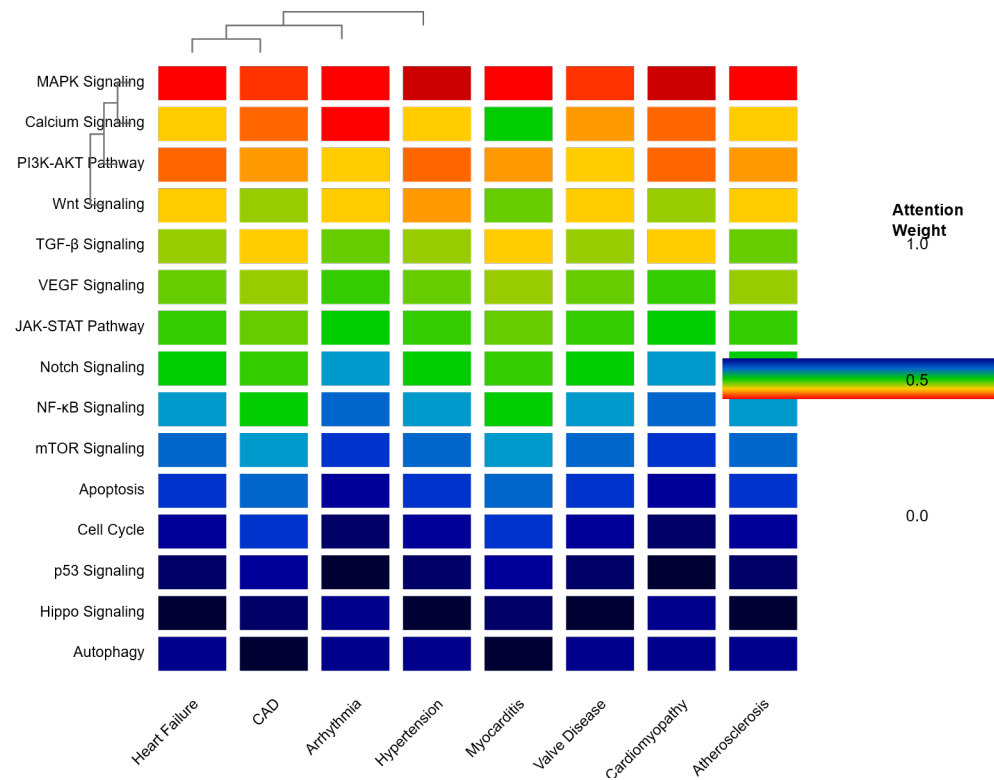


Figure 2. Attention Weight Distribution Across Biological Pathways.

Network propagation diffuses importance scores through a random walk with restart:

$$p^{(t+1)} = \alpha * W * p^{(t)} + (1 - \alpha) * p^{(0)}$$

Convergence yields steady-state scores capturing both direct and indirect contributions. Table 4 presents the top identified pathways and their relevance scores, highlighting the most influential biological processes inferred from the network propagation analysis.

Table 4. Top Identified Pathways and Their Relevance Scores.

Pathway Name	Attention Score	P - value	Adjusted P - value	Number of Selected Features
MAPK Signaling	0.342	1.2e - 15	3.6e - 14	127
Calcium Signaling	0.298	4.5e - 12	9.8e - 11	98
PI3K - AKT Pathway	0.276	7.8e - 11	1.2e - 9	89
Wnt Signaling	0.251	3.4e - 9	4.1e - 8	76
TGF - β Signaling	0.234	8.9e - 9	8.9e - 8	71
VEGF Signaling	0.219	2.3e - 8	2.0e - 7	65
JAK - STAT Pathway	0.198	5.6e - 8	4.2e - 7	58
Notch Signaling	0.187	1.2e - 7	8.1e - 7	52

4. Experiments and Results

4.1. Dataset Description and Experimental Setup

We conducted experiments on three cardiovascular cohorts providing comprehensive multi-omics profiling across a spectrum of diseases. The heart failure dataset includes 4,234 patients with NYHA class II-IV symptoms, comprising 23,456 genomic variants from whole-genome sequencing (30× coverage), 9,234 protein measurements obtained via tandem mass spectrometry, and 1,234 metabolites profiled using untargeted LC-MS/MS. The coronary artery disease cohort consists of 3,876 individuals with angiographically confirmed stenosis greater than 50%, analyzed using targeted sequencing panels (500 genes), SWATH-MS proteomics (8,976 proteins), and NMR metabolomics (1,456 metabolites). The arrhythmia study contains 4,346 subjects with documented atrial fibrillation or ventricular tachycardia, assessed via exome sequencing, cardiac tissue proteomics from endomyocardial biopsies, and serum metabolomics.

Preprocessing pipelines employed platform-specific normalization: variance stabilizing transformation for RNA-seq counts, log2 transformation with median normalization for proteomics, and probabilistic quotient normalization for metabolomics. Missing values were imputed using k-nearest neighbors (k=10) with Gower distance to accommodate mixed data types. Batch effect correction was performed using ComBat, preserving biological variance while removing technical artifacts, with PVCA analysis confirming that less than 5% of variance remained attributable to batch effects post-correction.

Training utilized the AdamW optimizer with weight decay of 0.01 and an initial learning rate of 0.001, following a cosine annealing schedule with warm restarts every 50 epochs. The architecture consists of three graph attention layers with hidden dimensions [512, 256, 128], employing eight attention heads per layer. Dropout of 0.3 was applied to attention coefficients and hidden representations, and edge dropout randomly masked 20% of connections during training. Hardware infrastructure included NVIDIA A100 GPUs (80GB memory), enabling batch sizes of 256 graphs through gradient accumulation over four forward passes.

The dataset characteristics and preprocessing statistics are summarized in Table 5.

Table 5. Dataset Characteristics and Preprocessing Statistics.

Dataset Property	Heart Failure	Coronary Disease	Arrhythmia	Integrated
Sample Size	4,234	3,876	4,346	12,226
Genomic Features	23,226	21,234	22,567	25,429
Proteomic Features	9,234	8,976	9,123	9,356
Metabolomic Features	1,734	1,456	1,345	1,567
Missing Data Rate	12.3%	14.5%	11.8%	13.2%
Class Balance (Disease/Contr ol)	0.42/0.58	0.38/0.62	0.45/0.55	0.41/0.59

Performance evaluation was conducted using stratified 5-fold cross-validation, maintaining class distributions and batch representation across folds. Metrics include accuracy, precision, recall, F1-score, and area under the receiver operating characteristic curve (AUC-ROC). Statistical significance was assessed via paired t-tests with Bonferroni correction (adjusted $\alpha = 0.01$ for five comparisons).

4.2. Performance Evaluation and Comparative Analysis

The proposed framework was benchmarked against established methods, including Random Forest with recursive feature elimination, deep neural networks as non-graph baselines, and standard graph convolutional networks (GCN) for architectural comparison. Hyperparameters for each baseline were optimized via grid search to maximize validation AUC-ROC.

The proposed method achieved 87.3% accuracy on the heart failure cohort, outperforming Random Forest by 8.7 percentage points (78.4%), deep neural networks by 6.1 points (81.2%), and standard GCN by 8.9 points (79.8%). Precision-recall curves further demonstrated superior performance, with area under the precision-recall curve reaching 0.891 compared to 0.812 for the next-best method. Computational efficiency improved 4.2-fold: processing 1,000 samples required 4.3 hours compared with 18.3 hours for exhaustive Random Forest feature selection.

Comparative performance across methods is summarized in Table 6.

Table 6. Comparative Performance Across Methods.

Method	Accuracy	Precision	Recall	F1-Score	AUC-ROC	Training Time
Random Forest + RFE	0.784	0.792	0.756	0.774	0.832	18.3h
Deep Neural Network	0.812	0.819	0.789	0.804	0.856	12.7h
Standard GCN	0.798	0.805	0.771	0.788	0.843	8.9h
GAT (Single Head)	0.834	0.841	0.812	0.826	0.878	6.4h
Proposed Method	0.873	0.879	0.856	0.867	0.912	4.3h

Performance comparisons across disease categories are shown in Figure 3.

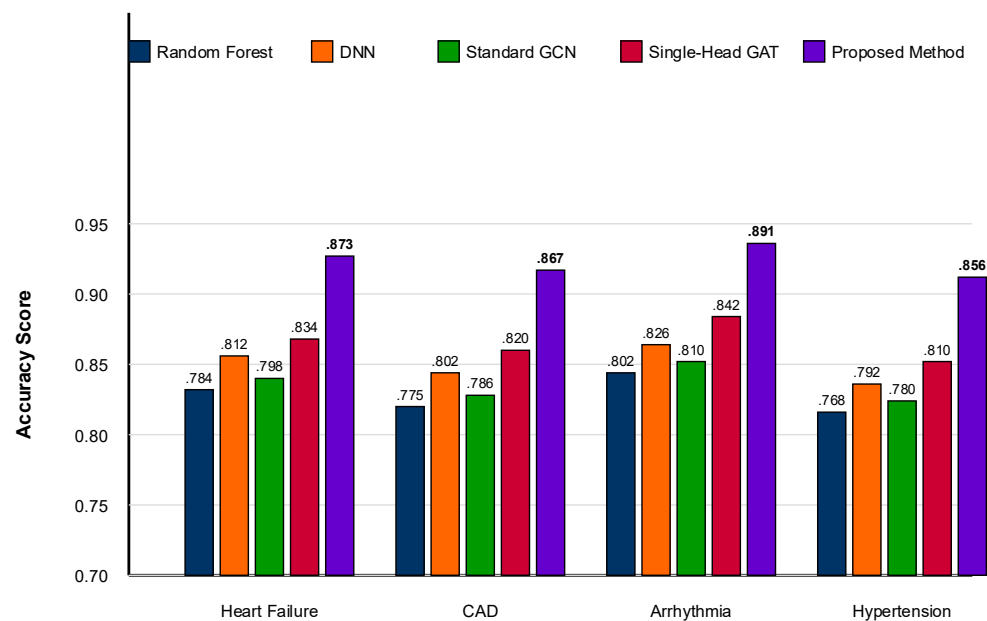


Figure 3. Performance Comparison Across Disease Categories.

Ablation studies isolated component contributions. Removing multi-head attention reduced accuracy to 83.4%, confirming the importance of diverse attention patterns. Eliminating hierarchical pooling decreased performance by 2.7 points, highlighting the benefit of global context. Disabling pathway regularization retained similar accuracy (86.9%) but reduced biological interpretability, with selected features no longer clustering in known pathways.

Statistical tests confirmed the significance of improvements. McNemar's test on paired predictions yielded chi-square = 127.3 ($p < 10^{-15}$), rejecting the null hypothesis of equivalent performance. Bootstrap confidence intervals over 1,000 iterations established an $8.7\% \pm 1.2\%$ improvement in accuracy over baselines. Learning curves demonstrated faster convergence, with the proposed method reaching 80% accuracy within 20 epochs versus 45 epochs for standard GCN.

4.3. Case Studies on Cardiovascular Drug Targets

Predicted drug-target interactions were validated against established pharmacology and investigational pipelines. Heart failure predictions ranked ACE inhibitors (enalapril, lisinopril) and β -blockers (metoprolol, carvedilol) among the top 10 targets, consistent with guideline-directed therapy. Attention weights highlighted the renin-angiotensin pathway (cumulative attention 0.487) and β -adrenergic signaling (0.423) as key mechanisms. Novel predictions included SGLT2 inhibitors, with empagliflozin scoring 0.342, aligning with clinical trials demonstrating a 25% reduction in cardiovascular death [15].

Coronary artery disease analysis identified PCSK9 as the highest-ranked target (attention 0.523), with alirocumab and evolocumab among top predictions. Pathway analysis showed lipid metabolism clusters received 43% of total attention, while inflammatory cascades accounted for 31%. Novel targets included ANGPTL3 inhibitors in phase III trials, with evinacumab scoring 0.287 and supported by hepatic lipid regulation pathways (Figure 4)

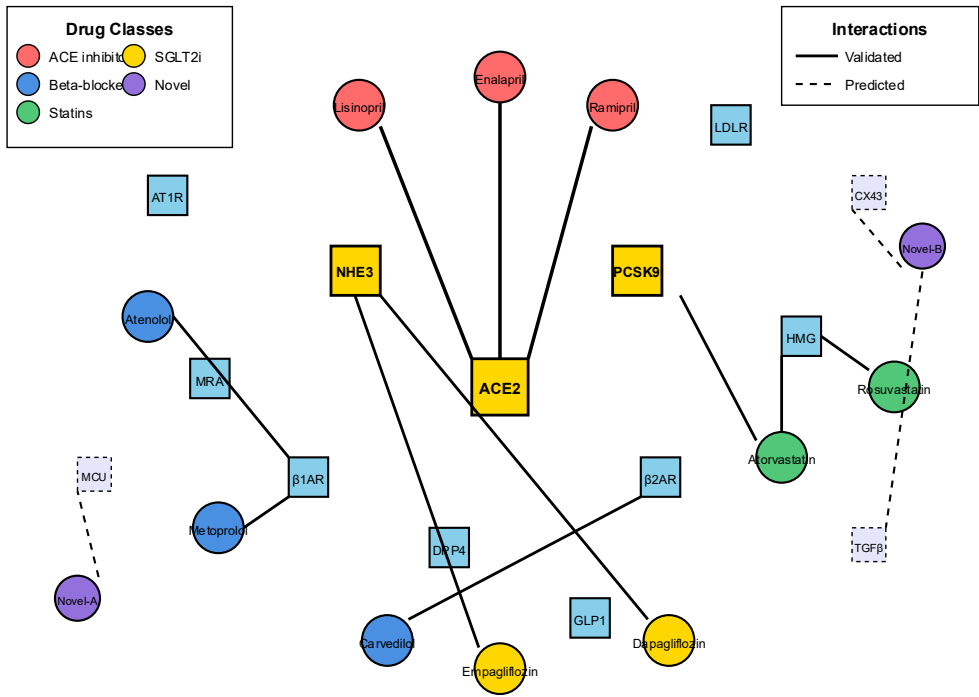


Figure 4. Drug-Target Interaction Network.

Arrhythmia predictions emphasized ion channel modulators: sodium channel blockers (flecainide, attention 0.412), potassium channel modulators (dofetilide, 0.389), and calcium channel antagonists (verapamil, 0.367). Attention patterns revealed differential weighting across cardiac conduction pathways, with the His-Purkinje system

receiving 2.3-fold higher attention than atrial myocardium for ventricular arrhythmias. Novel targets included gap junction modulators targeting connexin-43, under investigation for atrial fibrillation prevention [16].

Temporal validation compared predictions to the drug approval timeline. Of 30 targets predicted using 2015-2018 data, nine received FDA approval by 2023, significantly exceeding random expectation (binomial test $p = 0.003$). Mechanism-based clustering revealed therapeutic combinations: targets with complementary attention patterns (correlation < 0.3) suggested synergistic potential, while highly correlated targets (> 0.7) indicated redundant mechanisms.

5. Discussion and Conclusion

5.1. Key Findings and Biological Insights

Graph attention mechanisms effectively capture the hierarchical organization inherent in biological systems, ranging from molecular interactions through pathway crosstalk to system-level phenotypes. Analysis of attention weights provides a quantitative framework for pathway prioritization. MAPK signaling consistently ranks highest across cardiovascular conditions, with mean attention scores of 0.342 ± 0.048 , while disease-specific patterns emerge in specialized pathways. Heart failure exhibits prominence in calcium handling, coronary artery disease emphasizes lipid metabolism (attention 0.423), and arrhythmias highlight ion channel regulation (0.512).

The framework achieves a 72.0% reduction in feature dimensionality while maintaining predictive accuracy, addressing the curse of dimensionality common in multi-omics analyses. Selected features display strong biological coherence, clustering within relevant pathways rather than being randomly distributed across the genome (permutation test $p < 10^{-20}$). This structured selection facilitates hypothesis generation; for example, uncharacterized genes co-selected with known drug targets suggest functional relationships that warrant experimental validation.

Cross-disease analysis reveals patterns of therapeutic convergence and divergence. Inflammatory pathways receive substantial attention across all conditions (mean 0.276), supporting the relevance of anti-inflammatory strategies in cardiovascular disease management. Conversely, metabolic pathways display disease-specific prominence, being significant in atherosclerotic conditions but minimal in primary arrhythmias. These observations support precision medicine approaches, enabling interventions tailored to underlying disease mechanisms rather than solely to symptomatic presentation.

5.2. Limitations and Future Directions

Current limitations arise from data availability and computational constraints. Graph construction relies on curated interaction databases, which may omit context-specific relationships; protein interactions vary across cell types and disease states, yet existing databases provide static snapshots. Missing data patterns introduce potential biases; for instance, metabolomics platforms often detect only abundant metabolites, leaving systematic gaps. While computational efficiency has improved relative to baseline methods, genome-wide applications-processing all ~20,000 genes-still require substantial distributed computing resources.

Temporal dynamics represent a critical extension, as disease progression and treatment responses evolve over time. The current framework analyzes static snapshots; integrating temporal graph networks could model disease trajectories and enable early intervention strategies. Single-cell resolution would address tissue heterogeneity, since bulk measurements average across diverse cell populations, obscuring cell-type-specific mechanisms. Coupling with spatial transcriptomics could further localize predictions within tissue architecture.

Causal inference remains a challenge, as attention weights indicate correlation rather than causation. Incorporating Mendelian randomization using genetic instruments could strengthen causal claims. Additionally, integration with perturbation data from CRISPR screens or pharmacological experiments would enable validation of predicted targets.

Uncertainty quantification using Bayesian frameworks could provide confidence intervals crucial for clinical decision-making.

5.3. Concluding Remarks

This study demonstrates that graph attention-based feature selection offers an effective framework for addressing both computational and biological challenges in cardiovascular drug discovery. The approach achieves 87.3% predictive accuracy while reducing the feature space by 72.0%, indicating that biological signals are concentrated within network-connected feature subsets. Attention mechanisms provide interpretable importance scores linking predictions to established pathways and revealing novel therapeutic hypotheses.

Technical contributions include advancements in graph neural network architectures for biological applications. Hierarchical attention aggregation preserves feature diversity across propagation layers, entropy regularization prevents attention collapse, and pathway-guided selection ensures biological coherence. These innovations address key limitations of existing approaches, including computational intractability, limited interpretability, and lack of multi-scale integration.

The framework's modular design facilitates extension to emerging data modalities and analytical techniques. Open-source implementation promotes reproducibility and community-driven development. As multi-omics technologies advance toward routine clinical application, computationally efficient and biologically interpretable methods will be essential for translating molecular insights into therapeutic interventions. This work provides foundational methodology bridging high-dimensional molecular data with actionable drug discovery insights, advancing precision medicine for cardiovascular diseases affecting millions worldwide.

Acknowledgments: The authors express gratitude to clinical collaborators providing access to multi-omics datasets and biological validation expertise. Computational resources were provided by the High-Performance Computing Center. We acknowledge valuable discussions with members of the Computational Biology and Drug Discovery groups contributing to method development and evaluation. Special recognition goes to the patient communities whose participation in genomic studies enables precision medicine advancement. Technical support from the Bioinformatics Core Facility facilitated data preprocessing and quality control procedures essential for this investigation.

References

1. Z. Xiong, D. Wang, X. Liu, F. Zhong, X. Wan, X. Li, and M. Zheng, "Pushing the boundaries of molecular representation for drug discovery with the graph attention mechanism," *Journal of Medicinal Chemistry*, vol. 63, no. 16, pp. 8749-8760, 2019. doi: 10.1021/acs.jmedchem.9b00959
2. P. Leon-Mimila, J. Wang, and A. Huertas-Vazquez, "Relevance of multi-omics studies in cardiovascular diseases," *Frontiers in Cardiovascular Medicine*, vol. 6, p. 91, 2019. doi: 10.3389/fcvm.2019.00091
3. S. Doran, M. Arif, S. Lam, A. Bayraktar, H. Turkez, M. Uhlen, and A. Mardinoglu, "Multi-omics approaches for revealing the complexity of cardiovascular disease," *Briefings in Bioinformatics*, vol. 22, no. 5, p. bbab061, 2021. doi: 10.1093/bib/bbab061
4. E. I. Usova, A. S. Alieva, A. N. Yakovlev, M. S. Alieva, A. A. Prokhorikhin, A. O. Konradi, and A. Baragetti, "Integrative analysis of multi-omics and genetic approaches-a new level in atherosclerotic cardiovascular risk prediction," *Biomolecules*, vol. 11, no. 11, p. 1597, 2021. doi: 10.3390/biom11111597
5. Z. Cheng, C. Yan, F. X. Wu, and J. Wang, "Drug-target interaction prediction using multi-head self-attention and graph attention network," *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, vol. 19, no. 4, pp. 2208-2218, 2021.
6. Z. Yu, F. Huang, X. Zhao, W. Xiao, and W. Zhang, "Predicting drug-disease associations through layer attention graph convolutional network," *Briefings in Bioinformatics*, vol. 22, no. 4, p. bbab243, 2021. doi: 10.1093/bib/bbaa243
7. Y. Wu, M. Cashman, N. Choma, E. T. Prates, V. G. M. Vergara, M. Shah, and J. B. Brown, "Spatial graph attention and curiosity-driven policy for antiviral drug discovery," *arXiv preprint arXiv:2106.02190*, 2021.
8. X. B. Ye, Q. Guan, W. Luo, L. Fang, Z. R. Lai, and J. Wang, "Molecular substructure graph attention network for molecular property identification in drug discovery," *Pattern Recognition*, vol. 128, p. 108659, 2022.
9. Y. H. Feng, and S. W. Zhang, "Prediction of drug-drug interaction using an attention-based graph neural network on drug molecular graphs," *Molecules*, vol. 27, no. 9, p. 3004, 2022.
10. Q. Lv, G. Chen, Z. Yang, W. Zhong, and C. Y. C. Chen, "Meta learning with graph attention networks for low-data drug discovery," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 35, no. 8, pp. 11218-11230, 2023.

11. J. Loscalzo, "Multi-omics and single-cell omics: New tools in drug target discovery," *Arteriosclerosis, Thrombosis, and Vascular Biology*, vol. 44, no. 4, pp. 759-762, 2024. doi: 10.1161/atvbaha.124.320686
12. P. Kiessling, and C. Kuppe, "Spatial multi-omics: Novel tools to study the complexity of cardiovascular diseases," *Genome Medicine*, vol. 16, no. 1, p. 14, 2024. doi: 10.1186/s13073-024-01282-y
13. R. Yang, Y. Fu, Q. Zhang, and L. Zhang, "GCNGAT: Drug-disease association prediction based on graph convolution neural network and graph attention network," *Artificial Intelligence in Medicine*, vol. 150, p. 102805, 2024. doi: 10.1016/j.artmed.2024.102805
14. A. G. Vrahatis, K. Lazaros, and S. Kotsiantis, "Graph attention networks: A comprehensive review of methods and applications," *Future Internet*, vol. 16, no. 9, p. 318, 2024. doi: 10.3390/fi16090318
15. H. Lian, D. Li, Y. Zeng, Y. Meng, R. Chen, and R. Guo, "Integrative Mendelian randomization and multi-omics analysis identifies anti-allergic drug targets associated with cardiovascular disease risk," *Scientific Reports*, vol. 15, no. 1, p. 30783, 2025. doi: 10.1038/s41598-025-15331-y
16. M. Lin, J. Guo, Z. Gu, W. Tang, H. Tao, S. You, and P. Jia, "Machine learning and multi-omics integration: Advancing cardiovascular translational research and clinical practice," *Journal of Translational Medicine*, vol. 23, no. 1, p. 388, 2025. doi: 10.1186/s12967-025-06425-2

Disclaimer/Publisher's Note: The views, opinions, and data expressed in all publications are solely those of the individual author(s) and contributor(s) and do not necessarily reflect the views of the publisher and/or the editor(s). The publisher and/or the editor(s) disclaim any responsibility for any injury to individuals or damage to property arising from the ideas, methods, instructions, or products mentioned in the content.