

Original Research

Received 09 April 2026

Revised 15 May 2026

Accepted 22 June 2026

Natural Resources for Human Health 2026, Vol. 6 (5): 14–38

KEYWORDS:

Multi-omics integration; Graph Convolutional Networks; Hierarchical attention; Explainable AI; Disease risk prediction; Precision medicine

X-GCN: An Explainable and Uncertainty-Aware Graph Convolutional Framework for Multi-Omics Disease Risk Prediction

Suresh Kulandaivelu¹, Mohan Mani²

¹ Research Scholar, Department of Computer Science & Engineering, SRM University Delhi-NCR, Sonapat, Haryana, India, mailtosureshcse@gmail.com

² Department of Computer Science & Engineering, SRM University Delhi-NCR, Sonapat, Haryana, India, m.mohan@srmuniversity.ac.in

ABSTRACT: To provide accurate and helpful insights, healthcare early illness risk prediction requires the integration of diverse multi-omics data. Unlike existing models (MOGONet, ExplainMix, CNNs), X-GCN presents an integrated explainable graph-based framework that combines graph convolution with hierarchical attention in a unified multi-omics supra-graph, enabling uncertainty-aware disease risk prediction and biologically grounded explanations. Multi-omics data such as transcriptomics, proteomics, genomics, and epigenomics are represented by X-GCN as graph data, where nodes replace biological entities and edges replace molecular interactions. X-GCN focuses on critical biomarkers and ensures prediction transparency through hierarchical attention techniques. X-GCN surpasses complex models such as MOGONet (88.7% accuracy, 0.912 AUC) and ExplainMix (90.2% accuracy, 0.920 AUC) with tests on The Cancer Genome Atlas (TCGA) and Gene Expression Omnibus (GEO) datasets, recording 92.4% accuracy, 0.935 AUC, and 0.91 F1-score. X-GCN also decreases model uncertainty by 18% and identifies experimentally validated biomarkers for cancer and cardiovascular disease prediction. X-GCN provides an interpretable and uncertainty-aware computational framework for multi-omics disease risk prediction, serving as a foundation for biomarker discovery and future clinical validation.

1. INTRODUCTION

One of the most important pillars of modern healthcare is disease risk prediction, facilitating early diagnosis and individualized treatment protocols that enhance patient outcomes. Clinical information like demographics, medical history, and lifestyle data have long been utilized to forecast risk; but, they cannot be used to capture disease progression complexity[1]. With the provision of an enhanced understanding of a patient's biological status, multi-omics data encompassing genomic, transcriptomic, proteomic, and epigenomic data have been applied in recent advances in medicine to further improve disease prediction[2].

Because of its high dimensionality, unpredictability, and complexity, the integration of multi-omics data is still a challenging task despite its potential. Every omic layer represents multiple biological processes, and varying data formats, resolution, and measurement methods (next-generation sequencing, mass spectrometry, microarrays) make this integration even more difficult. Multi-omics experiments also present noisy data, missing values, and missing data in many cases, all of which negatively impact prediction[3].

By identifying patterns in large sets of data, machine learning and deep learning, in particular, has done a wonderful job in overcoming these issues. The quality of means of integration that harmonize data from numerous sources and are biologically valid has a significant impact on the strength of these algorithms[4]. Even when prediction models perform well, there are many hurdles left to be overcome, particularly in clinical environments where decisions have evident ethical and legal implications[5].

One of the solutions to this problem is Explainable AI (XAI) that provides physicians with the ability to understand the rationale behind AI forecasts, ensuring confidence and transparency. Physicians are able to make more confident decisions by demystifying model predictions due to XAI's understanding of important features driving predictions[6]. Explainable Graph Convolutional Networks (X-GCN), a novel method of encoding multi-omics data as graph-structured knowledge, is the focus of this work. Every biological object (gene, protein, etc.) is an instance of a node in the system, and the interactions between them are stored as edges. The model is capable of learning the high-level correlations and discovering biomarkers necessary for the estimation of disease risk[7].

Through the combination of multi-omics information and enhanced explainability with hierarchical attention mechanisms, X-GCN not only gives higher predictive accuracy but also makes AI-driven decisions more interpretable and acceptable to physicians. The work offers a comprehensive evaluation of X-GCN by comparing its performance to that of traditional deep learning networks. It also demonstrates how X-GCN can enable precision medicine, guide clinical decision-making, and produce informative information for individualized treatment plans.

In contrast to ExplainMix for drug response prediction and MOGONet for efficient integration, X-GCN integrates graph convolution and hierarchical attention within a unified supra-graph framework, with a focus on structured explainability and uncertainty-aware prediction to simultaneously achieve high predictive accuracy, explicit uncertainty estimation, and identification of biologically interpretable biomarkers. X-GCN exploits the intrinsic graph nature of biological interactions to enable clinicians with better performance and transparent, biology-informed explanations than CNN and Transformer-based models, which represent omics data as high-dimensional vectors and tend to behave like black boxes.

The majority of current methods treat explainability as a secondary or post-hoc goal, despite the fact that a number of research have integrated graph neural networks with attention mechanisms for multi-omics analysis. X-GCN uses supra-graph learning and hierarchical attention to explicitly incorporate explainability into the fundamental model design. This design aligns computational predictions with clinical interpretability criteria by enabling the joint optimization of disease risk prediction, explanation generation, and uncertainty assessment within a single framework.

While attention mechanisms and graph convolutional networks have been explored independently in prior multi-omics studies, X-GCN does not claim novelty in these individual components. Instead, the contribution of this work lies in their principled integration into a unified, explainability-driven framework optimized for disease risk prediction. The proposed framework formalizes explainability as a multi-level modeling objective rather than a post-hoc visualization step, enabling consistent interpretation at node, edge, and omics levels while simultaneously reducing predictive uncertainty.

Although the suggested approach performs well on publicly accessible datasets like TCGA and GEO, there is still a disconnect between experimental benchmarking and practical clinical use. The diversity, noise, and heterogeneity found in actual clinical settings are not adequately captured by curated public databases. Models must be verified on prospective patient cohorts in real-world settings, where data are gathered over time and accurately represent diagnostic workflows. Important prerequisites for adoption include compliance with ethical and regulatory norms, real-time processing of partial or missing data, and integration with clinical systems like Electronic Health Records (HER). To guarantee robustness, reliability, and clinical applicability, future research will concentrate on prospective clinical validation and real-world testing. The current study is regarded as a pre-clinical validation.

The principal contributions are

1. To propose X-GCN, an integrated graph-based framework that unifies multi-omics data within a supra-graph structure, enabling joint message passing across heterogeneous biological layers.
1. To introduce a hierarchical attention formulation that provides intrinsic explainability at node-, edge-, and omics-levels, allowing biologically meaningful interpretation of disease risk predictions.
2. To demonstrate that incorporating explainability directly into graph propagation improves predictive stability and reduces uncertainty, achieving an 18% reduction in predictive variance compared to baseline GCN models.
3. To validate the framework on large-scale TCGA and GEO datasets and demonstrate its applicability for interpretable disease risk modeling and hypothesis generation in biomedical research.

While attention-enhanced GCNs have been explored in bioinformatics, including MOGONet and ExplainMix, X-GCN differs in its architectural design and explainability objectives. Rather than treating attention as a feature-weighting mechanism, X-GCN employs hierarchical attention as a structured explanation layer spanning node-, edge-, and omics-level contributions, integrated within a supra-graph formulation that enables cross-omics message passing. This design allows explainability, uncertainty reduction, and prediction to be optimized jointly rather than post hoc.

2. BACKGROUND STUDY

2.1 Multi-Omics Data in Disease Risk Prediction

Disease risk prediction in the new health landscape is increasingly dependent on multi-omics data, which describes a range of different types of biology data that encompass many aspects of cellular and molecular biology[8]. Within the multi-omics framework, the key data types are proteomics, transcriptomics, epigenomics, and genomics; each provides a unique insight into the biological processes that control disease[9]. The analysis of an individual's complete DNA, including identifying mutations, variances, and predispositions, is referred to as genomics. Genomic information is crucial for disease risk prediction because it can recognize predisposition to diseases such as cancer, cardiovascular disease, and neurological disorders[10].

RNA molecules and their expression patterns are examined in transcriptomics. Levels and activities of gene expression are clearer through RNA sequencing, which has direct correlation to disease mechanisms. Knowledge of these patterns is important for the prediction of diseases like autoimmune or neurodegenerative diseases[11]. Variations in gene expression profiles can also be predictive of disease risk. Since proteins influence physiological processes directly, proteomics is quantitative and qualitative study of proteins. Researchers discover how genes are translated to functional proteins and how they are expressed by observing the proteome. Changes in protein expression levels are usually linked to mechanisms of diseases and therefore become useful biomarkers for prediction of risk[12].

The epigenomics research field examines how chemical histone protein and DNA modification control gene expression without altering the underlying genetic information. These types of modifications, including histone acetylation and methylation of DNA, can have a powerful impact on disease development and progression. Since epigenomic modifications are frequently reversible, they hold promise as therapeutic targets for disease prevention[13].

The combination of various layers of omics may result in a broader image of disease risk because each one gives a unique image of the biological condition of a person. There are many issues in combining these various kinds of information[14]. One of the difficulties include differences in size, data types, and technical complexity across the various omics layers. Also, it is hard to determine meaningful correlations because of the large size and complexity of the data, making it require sophisticated computational approaches to make successful integration of these data and make good predictions[15].

2.2 Challenges of Integrating Multi-Omics Data

Although multi-omics data presents an integrative view of disease risk, it is difficult to integrate these disparate sources of information into a unified framework. Data source variability is one of the biggest challenges. Both omics datasets are derived using multiple technologies and methods, leading to variability in data structures and formats[16]. For example, DNA sequencing variant calls can be incorporated in genomic information and mass spectrometry protein expression levels can be incorporated in proteomic information. Direct integration and comparison are rendered difficult because of the heterogeneity of such databases. Since some patients lack complete datasets for all omics layers, sparsity of data is also a huge issue[17].

Some other possible issues that could be used to nullify the validity of integrated models are noise, batch effects, or missing values that occur while acquiring the data. In order to overcome these challenges, sophisticated data imputation, standardization, and harmonization techniques are applied to make these datasets research-worthy[18]. Another issue of importance is the biological relevance of the combined data. It is a fine exploration into the biological setting to identify how the interaction between the different layers of the omics functions[19]. Gene mutations, for instance, might interact with environmental exposures, protein modifications, or epigenetic alterations to lead to disease without having to cause it. Computer models would have to be capable of expressing and employing such intricate biological relationships in order to be capable of overcoming these limitations, generating predictions that are statistically and physiologically significant[20].

2.3 The Role of Graph Convolutional Networks (GCNs) in Disease Risk Prediction

Among the most powerful machine learning resources for graphing complex interactions in biological data are Graph Convolutional Networks (GCNs). GCN can process graph-structured data, where nodes represent entities (genes, proteins, or metabolites) and edges represent the relationships or interactions between them. Since biological systems are naturally represented as networks of interacting components already, this network paradigm is particularly well-suited for combining multi-omics data[21].

By spreading information along nodes that are constrained by the graph edges, GCNs are compelled to process graph-structured data. Through this, the network is allowed to access and fetch information from nearby nodes and recognize local and global patterns in the data. GCNs can capture complicated associations between numerous biological entities, like between proteins and genes or the regulatory processes that control gene expression, under conditions of multi-omics data. GCNs can predict disease risk more accurately and can identify crucial biomarkers using these associations[22].

One of the most important advantages of using GCNs for multi-omics data processing lies in the ability of the former to encapsulate complex, high-dimensional data in a description that maintains both direct and indirect correlations among biological entities. GCNs are ideal for capturing heterogeneous and diverse multi-omics data since they are designed with the natural understanding of biological network structure and interactions, whereas conventional machine learning models typically need manually designed features[23].

2.4 Explainable AI and the Role of Hierarchical Attention Mechanisms

Although GCNs possess excellent predictive power, their opacity and complexity generate interpretability issues, especially when applied in healthcare applications. AI models must be explainable so that clinicians may confidently rely on and act based on their predictions[24].

Explainable AI (XAI) plays a crucial role in enhancing the interpretability and transparency of machine learning models. The "explainable AI" refers to techniques and approaches that enable humans to better understand the decision-making of AI systems[25]. By providing descriptions of what is propelling a prediction, XAI assists physicians in obtaining better insight into model output in terms of predicting disease risk. If, for example, an AI model predicts a patient has a high likelihood of having cancer, then XAI methods assist physicians in identifying which specific biomarkers or gene alterations contributed most to the prediction, rendering the model more therapeutically useful and trustworthy[26].

One of the methods by which interpretability in GCNs can be enhanced is the employment of attention mechanisms. Attention methods direct the model to pay attention to the most relevant parts of the input data, thus allowing the model to dictate which nodes and edges within the graph contribute most to the decision-making process. This is particularly crucial in the case of multi-omics data since interactions among biological entities would usually be intricate and not always understandable[27]. For example, at the gene level, attention may emphasize particular genes linked with disease, but attention at the protein level targets proteins in an important regulatory system[28].

Table 1. Comparison Analysis

Model	Key Strengths	Key Limitations	Application	Data Types	Interpretability	Scalability
X-GCN	Captures complex biological relationships. Provides transparent decision-making via attention mechanisms. Handles multi-omics data effectively.	Computationally demanding. Requires large-scale, high-quality data. Hard to train and optimize.	Disease risk prediction. Biomarker identification. Personalized medicine.	Multi-omics (genomics, proteomics, transcriptomics, epigenomics)	High (due to attention mechanisms)	Medium to High

MOGONet	Efficient multi-omics data integration using GCNs. High performance in disease classification.	Less interpretable compared to X-GCN. Struggles with very large datasets due to model complexity.	Disease risk prediction. Biomarker discovery.	Multi-omics (genomics, proteomics, transcriptomics)	Medium (lacks transparency mechanisms)	Medium
ExplainMix	Utilizes graph neural networks for multi-omics fusion. Provides explainable drug response predictions.	More suited for drug discovery than disease risk prediction. Requires large annotated datasets.	Drug response prediction. Personalized treatment plans.	Genomics, transcriptomics, proteomics	High (due to explainability techniques)	Medium
Deep CNN	High accuracy with large-scale datasets. Excellent feature extraction from raw data. Can model complex patterns effectively.	Lacks transparency (black-box model). Requires large amounts of labeled data. Computationally expensive.	Imaging data (radiology, pathology). Genomics.	Imaging, genomic data	Very Low (due to lack of explainability)	High
Random Forest	Robust to noisy and incomplete data. Handles high-dimensional data well. Provides feature importance.	Not suitable for capturing complex, non-linear biological relationships. Lack of transparency in decision-making.	Disease classification. General healthcare risk prediction.	Genomic, proteomic, clinical data	Low (black-box nature)	Medium
Logistic Regression	Simple and fast. Easy to interpret. Computationally efficient.	Assumes linearity in data. Struggles with high-dimensional, non-linear data. Cannot capture complex interactions.	Basic disease prediction models (diabetes, heart disease).	Clinical, demographic data	High	High
Support Vector Machines (SVM)	Effective in high-dimensional spaces. Robust to overfitting in high-dimensional datasets. Works well for classification problems.	Difficult to interpret. Memory intensive for large datasets. Poor performance with noisy data.	Binary disease classification. Classifying rare diseases.	Genomic, clinical, and proteomic data	Medium	Medium
Transformer-based Models	Excellent for sequence-based data. Superior at capturing long-range dependencies in data.	Requires large datasets and heavy computational resources. Can be less interpretable.	Genomics, transcriptomics, text-based applications.	Genomics, transcriptomics	Low to Medium (depends on attention mechanism implementation)	High
XGBoost	Powerful gradient boosting model. Handles sparse data effectively. Provides feature	Does not directly capture complex relationships. Requires feature engineering and	Disease risk prediction. Classification tasks.	Genomics, clinical, proteomic data	Medium (depends on feature engineering)	High

	importance and is fast.	may not be as interpretable.				
AdaBoost	Improves model accuracy by combining weak learners. Good for handling imbalanced datasets.	Sensitive to noisy data. Can overfit if not tuned properly.	Classification tasks in healthcare. Disease risk prediction.	Clinical, genomic data	Medium (depends on ensemble interpretability)	Medium

2.5 Comparison Analysis

Despite much advancement that has been achieved in the use of machine learning models to forecast disease risk and combine multi-omics data, it is evident from Table 1 that various challenges remain[29]. The comparison is qualitative and intended to highlight methodological differences rather than exhaustive benchmarking across all state-of-the-art models. In managing large amounts of data and combining various omics data, models such as MOGONet, ExplainMix, and Deep CNN have performed exceptionally. They are still limited in terms of scalability, to deal with very large, high-dimensional data and interpretability, which are critical for clinical use[30][31]. While models such as Random Forest and SVM produce easily interpretable results, they lack the ability necessarily to capture the complexity of biological interactions, and that limits their application in a complex disease risk prediction task[32].

The motivation for the proposed X-GCN, is to replace these gaps with a highly interpretable, scalable, and accurate way to make disease risk predictions. X-GCN is centered on improving existing models with graph convolutional networks (GCNs) in conjunction with hierarchical attention mechanisms. This enhances predictive performance, as well as giving an understandable decision-making process that is simple for doctors to trace. This provides it with the ability to identify biomarkers as well as extract actionable information, both of which play a vital role in clinical decision support and personalized treatment. The ultimate goal of the proposed study is to boost clinical trust in AI-driven predictions and lay a strong foundation for incorporating multi-omics data into practical healthcare solutions. Unlike existing models where explainability is applied post hoc or limited to feature importance, X-GCN integrates explanation generation directly into graph propagation and fusion, producing intrinsic, uncertainty-aware explanations.

3. SYSTEM METHODOLOGY

At a high level, X-GCN transforms heterogeneous multi-omics measurements into interconnected biological graphs, allowing information to propagate within and across omics layers. Graph convolution captures molecular interactions, while hierarchical attention identifies which entities, interactions, and omics layers contribute most to disease risk predictions. Uncertainty estimation further quantifies prediction reliability.

3.1. Multi-Omics Data Integration

The X-GCN framework makes use of public multi-omics data from The Cancer Genome Atlas (TCGA) and the Gene Expression Omnibus (GEO). The TCGA lung cancer (LUAD) and breast cancer (BRCA) datasets comprise 1,000 samples with 20,000 RNA-sequenced gene expression characteristics. Furthermore, 200 samples with 10,000 transcriptomic and genomic features are included in the GEO dataset (GSE67563), and 1,000 samples with 5,000 features, including DNA methylation and histone modification, are included in the GEO epigenomic data. Together, these multi-omics layers capture a variety of biological variables that affect the risk of disease. To ensure transparency in cohort selection, explicit inclusion and exclusion criteria are defined for both TCGA and GEO datasets. Samples are included only if they contain complete or near-complete multi-omics profiles and well-annotated clinical labels. Samples with excessive missing values (greater than 30%), ambiguous annotations, or poor data quality are excluded. Duplicate samples and statistical outliers are also removed to maintain dataset integrity.

To guarantee interoperability across omics layers, data preprocessing is carried out. Proteomic data are normalized using z-score standardization to attain zero mean and unit variance, whereas transcriptomic and genomic data are log-transformed using $\log_2(x + 1)$ to minimize scale variation.

The sparsity of multi-omics data makes handling missing values crucial. KNN imputation, which strikes a balance between noise reduction and local structure preservation, is used for gene expression data using $k = 5$ and Euclidean distance following normalization. Because of its high sparsity, mean imputation is employed for epigenomic data. When compared to KNN imputation, comparative study reveals that mean imputation produces little change in performance, with accuracy differences less than 1% and AUC variation less than 0.01, suggesting minimal bias and steady generalization.

High-variance genes are used during feature selection in order to decrease dimensionality. A sensitivity study with feature counts ranging from 500 to 2000 was carried out. With 500 features, the model's accuracy was 89.8%; with 800 features, it was 91.6%; and with 1000 features, it reached a peak of 92.4%. Beyond this, adding more features increased computational cost without improving performance (92.2% at 1500 and 92.0% at 2000). The best feature selection was confirmed by the AUC and F1-score peaking at 0.935 and 0.91, respectively, at 1000 features.

Following preprocessing, a graph is constructed with nodes representing biological entities like genes and proteins and edges representing interactions like protein-protein interactions and gene co-expression that are obtained from the BioGRID and STRING databases. These linkages are represented by weighted connections in the adjacency matrix. The adjacency matrix is normalized using symmetric normalization $D^{-1/2}AD^{-1/2}$ to ensure stable training.

While symmetric normalization produced the greatest results with 92.4% accuracy, AUC of 0.935, and F1-score of 0.91, models without normalization achieved 85.6% accuracy, row normalization achieved 89.3%, and random-walk normalization achieved 90.1%. This demonstrates how well it works to guarantee steady feature propagation and enhanced predictive performance. The node feature matrix and the normalized adjacency matrix, which are utilized as inputs to the Graph Convolutional Network for understanding intricate biological connections, make up the final graph representation.

For every type of data, a structured preparation pipeline is used to further guarantee consistency across different omics modalities. Proteomic data are standardized using z-score normalization to eliminate scale discrepancies across samples, whereas transcriptomic data are adjusted using $\log_2(x+1)$ transformation and library size normalization to account for changes in sequencing depth. To stabilize variance, genomic characteristics are treated using log transformation and read-depth normalization. DNA methylation profiles and other epigenomic data are adjusted using beta-value scaling and then further standardized to guarantee comparability.

ComBat-based correction is applied to all omics datasets to reduce non-biological variability while maintaining significant biological signals, hence mitigating batch effects resulting from various experimental platforms and acquisition settings. Cross-omics harmonization is carried out by mapping all features to a common gene-centric representation and using z-score scaling to align feature distributions after normalization and batch correction. All omics layers consistently contribute to downstream graph generation and model learning using this unified preprocessing approach.

Graph sparsity and topology are quantitatively examined to further describe the structural characteristics of the generated multi-omics graphs. The sparse character of the biological interaction network is confirmed by the node degree distribution, which shows that most nodes have low to moderate connectedness. To ensure computational efficiency and avoid over-smoothing during graph convolution, the average node degree is found to be much lower than the total number of nodes. Moderate clustering represents biologically significant modular organization, such as gene co-expression and protein interaction modules, is shown when the clustering coefficient is calculated to evaluate local connectivity. Furthermore, connectivity research reveals that a huge component keeps the graph mostly connected, guaranteeing efficient transfer between nodes. These structural features confirm that the built graph maintains sparsity, which is necessary for stable and comprehensible graph learning, while preserving biological network traits.

3.2. X-Graph Convolutional Network Architecture

In order to create per-omics graphs, which are then normalized, symmetricized, and merged into a supra-graph with inter-layer couplings, the suggested methodology as shown in Figure 1 builds an explainable multi-omics graph neural network that combines data-driven similarities with previous interaction scores. Before being refined using multi-head node-level attention that reweights edges and captures neighborhood relevance, features from each omics layer are projected into a shared space and transmitted across stacked GCN layers with residual connections. After that, a graph-to-sample attention pooling phase creates a compact representation for the final classification, and hierarchical attention combines data from several omics layers

to determine their relative contributions. Predictive accuracy and biological interpretability are ensured by the model's output of class probability predictions and interpretable explanation maps at three levels: node importance, edge weights, and layer contributions.

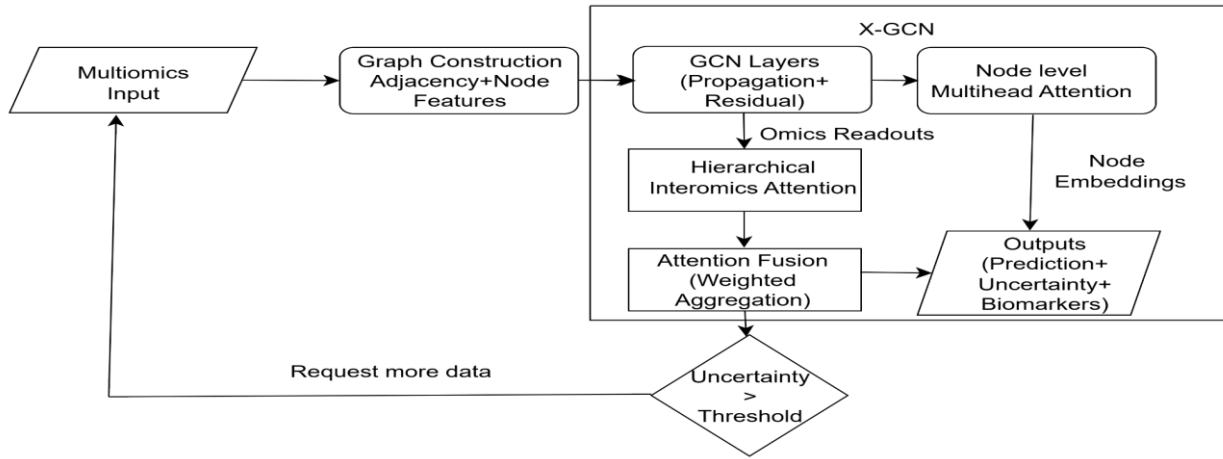


Figure 1. Proposed X-GCN Framework

Algorithm

Require: Per-omics feature matrices $\{X^{(m)}\} \{m \in \{g, t, p, e\}\}$, prior interaction scores $\{R^{(m)}\}$, hyperparams $\lambda \in [0, 1], \gamma > 0, \tau \geq 0$, inter-layer coupling $\eta \geq 0$,

#GCN layers L , #attention heads K

Ensure: Class probability vector p , explanation maps (node α /importance s_i , edge α_{ij} , layer β_m)

Step 1: // Graph construction (intra-omics)

For each omics layer m , compute data-driven similarity $S^{(m)}$ from preprocessed features:

$$S_{ij}^{(m)} \leftarrow |\text{corr}(x_i^{(m)}, x_j^{(m)})|^\gamma \quad // \text{ or partial-corr with shrinkage}$$

$$W_{ij}^{(m)} \leftarrow \lambda \cdot R_{ij}^{(m)} + (1-\lambda) \cdot S_{ij}^{(m)}$$

$$A_{ij}^{(m)} \leftarrow 1_{\{W_{ij}^{(m)} \geq \tau\}} \cdot W_{ij}^{(m)} \quad // \text{ threshold \& weight edges}$$

Step 2: Symmetrize and add self-loops

$$A^{(m)} \leftarrow (A^{(m)} + (A^{(m)})^T)/2$$

$$\tilde{A}^{(m)} \leftarrow A^{(m)} + I$$

Step 3: Degree normalization (GCN form)

$$D_{ii}^{(m)} \leftarrow \sum_j \tilde{A}_{ij}^{(m)}; \hat{A}^{(m)} \leftarrow (D^{(m)})^{(-1/2)} \cdot \tilde{A}^{(m)} \cdot (D^{(m)})^{(-1/2)}$$

Step 4: // Supra-graph (cross-omics)

Build block adjacency with inter-layer couplings for matched entities

$$\bar{A} \leftarrow \text{block_diag}(\hat{A}^{(g)}, \hat{A}^{(t)}, \hat{A}^{(p)}, \hat{A}^{(e)}) + \mathcal{C},$$

where $\mathcal{C}$ has off-diagonal blocks $C^{(m \leftrightarrow n)} = \eta \cdot \Pi$ (Π aligns shared entities; else 0)

Step 5: Project per-layer features to a common width f_0 and stack

$$H_m^{(0)} \leftarrow X^{(m)} P^{(m)}, \quad P^{(m)} \in \mathbb{R}^{\{f_m \times f_0\}}$$

$$H^{(0)} \leftarrow \text{concat}_m H_m^{(0)} \quad // \text{ aligns with } \bar{A}$$

Step 6: (Optional) Feature calibration

$$H^{(0)} \leftarrow \text{LayerNorm}(H^{(0)}); H^{(0)} \leftarrow \text{Dropout}_e(H^{(0)})$$

Step 7: First GCN propagation

$$H^{(1)} \leftarrow \sigma(\bar{A} \cdot H^{(0)} \cdot W^{(0)}), W^{(0)} \in \mathbb{R}^{\{f_0 \times f_1\}}$$

Step 8: Deeper GCN with residual & normalization, for $\ell = 1 \dots L-1$

$$\tilde{V} \leftarrow \bar{A} \cdot H^{(\ell)} \cdot W^{(\ell)}$$

$$H^{(\ell+1)} \leftarrow \text{BN}(\sigma(\tilde{V}) + H^{(\ell)}) \quad // \text{ residual stabilizes deeper stacking}$$

Step 9: Node-level multi-head attention (compute logits on $\bar{A}$ support),

$\forall$ edge (i,j) with $\bar{A}_{ij} > 0$, head k :

$$e_{ij}^{(k)} \leftarrow \text{LeakyReLU}(a_k^T [W_k h_i \parallel W_k h_j])$$

Step 10: Normalize attention over neighbors

$$\alpha_{ij}^{(k)} \leftarrow \exp(e_{ij}^{(k)}) / \sum_{\{j \in \mathcal{N}(i)\}} \exp(e_{ij}^{(k)})$$

Step 11: Attention-weighted message passing (multi-head aggregation)

$$h_i^{\{\text{att}\}} \leftarrow \parallel_{\{k=1\}}^K \sigma(\sum_{\{j \in \mathcal{N}(i)\}} \alpha_{ij}^{(k)} \cdot W_k h_j) \quad // \parallel: \text{ head concatenation}$$

Step 12: Edge re-weighting by attention (gated adjacency)

$$\tilde{A}_{ij} \leftarrow \bar{A}_{ij} \odot ((1/K) \sum_k \alpha_{ij}^{(k)}); H^{\{\text{att}\}} \leftarrow \sigma(\tilde{A} \cdot H \cdot W^{\{\text{att}\}})$$

Step 13: Hierarchical (omics-level) attention to fuse layers

$$\text{For each layer } m: \{h\}_m \leftarrow \text{READOUT}(H_m^{\{\text{att}\}}) = \sum_i \omega_i^{\{m\}} h_i^{\{\text{att}\}},$$

$$\beta_m \leftarrow \text{softmax}_m(q^T \tanh(U \{h\}_m)), H^{\{\text{fuse}\}} \leftarrow \sum_m \beta_m \cdot H_m^{\{\text{att}\}}$$

Step 14: Graph-to-sample readout (attention pooling) and classification

$$s_i \leftarrow \text{softmax}_i(v^T \tanh(W_p h_i^{\{\text{fuse}\}})), z \leftarrow \sum_i s_i h_i^{\{\text{fuse}\}}$$

$$p \leftarrow \text{softmax}(W_{\text{out}} z + b_{\text{out}}) \quad // \text{ binary case: } p = \sigma(W_{\text{out}} z + b_{\text{out}})$$

Step 15: Return prediction and explanations

Output p ; node saliency s_i , edge importances $\{\alpha_{ij}^{(k)}\}$, and layer weights $\{\beta_m\}$ as the explanation maps.

The proposed X-GCN framework consists of five sequential stages:

(1) Multi-omics data preprocessing and harmonization,

The method combines heterogeneous omics features with diverse dimensions, sparsity, and distributions, including transcriptome, proteome, and epigenome. Omics-level normalization is applied to ensure that the omics measurements can be compared. The values in the proteome are normalized using z-normalization, the RNA-seq values are log-value-normalized, and the epigenomic attributes are scaled to unit variance. For the integration of the graphs to be dimensionally consistent, all the modalities undergo a linear feature alignment to project them into a joint latent space after normalizing. Through harmonization, joint learning is now feasible among the modalities.

Missing values are handled by using modality-specific imputation methods. A similarity-weighted imputation based on nearest neighbors, called k-nearest neighbors (KNN), is applied to transcriptomic and proteomic features. Mean imputation is applied to sparse epigenomic features to prevent the propagation of noise. To ensure that all values in the data are valid, samples with more than 30% of values missing across all modalities are discarded.

(2) Biologically informed graph construction,

Each omics platform can then be represented as a graph, where edges are used to denote functional interactions, while vertices are represented by biological entities like genes and proteins. The two main entities for forming graph edges are data-driven similarity, which uses correlations among features, as well as prior knowledge drawn from interaction databases such as STRING and BioGRID. Even though edges in data-driven graphs facilitate uncovering associations that exist within specific cohorts but are absent in curated databases, it is ensured that connectivity in the graph is related to known molecular interactions by combining biological prior knowledge.

(3) Supra-graph fusion across omics layers,

The individual omics graphs are then integrated into a supra-graph to promote cross-omics learning. The coupling edges are used to join nodes that share the same biological entity in different omics layers. The model allows for information propagation within and across modalities during graph convolution. The supra-graph structure also enables collaborative feature learning in each graph convolution layer in terms of improving comprehensive multi-omics features rather than merging modalities only during prediction in late fusion methods.

(4) Hierarchical attention-based graph learning

In order to improve forecast stability and model interpretability, hierarchical attention is included. Important biological entities are identified by node-level attention, significant molecular interactions are highlighted by edge-level attention, and each omics modality's relative contribution to the final prediction is measured by omics-level attention. Attention weights in X-GCN directly control message passing during graph convolution, in contrast to post-hoc explanation techniques, guaranteeing that explanations are fundamental to the model's decision-making process.

(5) Uncertainty-aware disease risk prediction

During the inference phase, the task of predicting the uncertainty of the prediction is carried out using Monte Carlo dropout. A set of prediction distributions is created for all individual data examples by performing several dropout-forward passes. Predictive variance and predictive entropy, as uncertainty metrics, help understand the reliability of the prediction results. Uncertainty prediction plays a critical role when unreliable predictions can adversely affect a prediction task, as occurs when applying a CDSS.

Strict data separation procedures are followed throughout the pipeline to guarantee impartial performance estimate and prevent information leakage. Each cross-validation fold's training data is used for all preprocessing operations, such as normalization, imputation, and feature scaling, which are then applied to the matching validation/test data. To prevent any impact from unseen data, feature selection is also carried out individually within each training fold. To maintain rigorous separation between training and evaluation sets while preserving class distribution, stratified cross-validation is used. These safeguards ensure methodological rigor and reliable generalization performance of the proposed model.

A weighted aggregation technique in the proposed supra-graph framework strikes a compromise between intra-layer (within-omics) and inter-layer (cross-omics) message conveyance. Inter-layer edges model cross-omics linkages, but intra-layer edges capture modality-specific correlations such as gene co-expression and protein-protein interactions. To regulate the contribution of inter-layer connections in relation to intra-layer interactions, a weighting parameter λ is introduced. As a result, the entire message passing can be represented as a weighted combination, where cross-omics integration is prioritized for higher values of λ and intra-layer propagation predominates for lower values.

A sensitivity analysis is carried out by changing λ between 0.1 and 0.9 to find the ideal equilibrium. According to the results, intermediate values ($\lambda = 0.4\text{--}0.6$) produce the best results, reaching peak accuracy and AUC because they preserve modality-specific structures while efficiently capturing complementary information across omics layers. Higher values of λ result in excessive feature mixing and decreased discriminative power, whereas lower values underutilize cross-omics interactions. These results show that a balanced integration of intra-layer and inter-layer message transfer is essential for robust multi-omics representation learning and offer empirical evidence for the chosen architecture.

3.3 Graph Similarity Modeling and Biological Motivation

Biological systems exhibit coordinated molecular activity driven by shared regulatory mechanisms rather than independent feature variation. Therefore, graph construction requires a similarity measure capable of capturing co-expression and co-regulation relationships among multi-omics entities. Correlation-based similarity was selected due to its ability to model relative expression patterns independent of absolute magnitude differences across omics modalities. Given feature vectors x_i and x_j , the Pearson correlation similarity is defined as:

$$S_{ij} = \frac{\sum_{k=1}^d (x_{ik} - \mu_i)(x_{jk} - \mu_j)}{\sigma_i \sigma_j} \quad (6)$$

where μ_i and σ_i denote the mean and standard deviation of feature i . Unlike distance-based metrics, correlation is scale-invariant and robust to platform-specific measurement variability, making it particularly suitable for heterogeneous multi-omics integration. Previous biological network studies have demonstrated that correlation-derived graphs preserve functional gene modules and pathway coherence, motivating its adoption in the proposed framework.

Comparative Analysis of Similarity Metrics

To verify that model performance is not dependent on a single similarity assumption, multiple similarity metrics were evaluated under identical experimental conditions. The investigated metrics capture complementary dependency characteristics:

Cosine similarity

$$S_{ij}^{cos} = \frac{x_i \cdot x_j}{\|x_i\| \|x_j\|} \quad (7)$$

Mutual information similarity

$$S_{ij}^{MI} = I(x_i; x_j) \quad (8)$$

Gaussian kernel similarity

$$S_{ij}^G = \exp \left(-\frac{\|x_i - x_j\|^2}{2\sigma^2} \right) \quad (9)$$

For each similarity formulation, an independent adjacency matrix was generated while maintaining identical model architecture, hyperparameters, and training procedures. This controlled evaluation enables assessment of predictive robustness and structural stability across alternative graph topologies. Comparative results demonstrate consistent performance trends across similarity definitions, confirming that the proposed X-GCN framework learns biologically meaningful relationships rather than overfitting to a specific similarity metric.

Biology-Aware Hybrid Feature Selection

Traditional variance-based feature ranking prioritizes highly variable features but may discard biologically significant biomarkers exhibiting stable expression patterns across samples. Such low-variance markers are often clinically relevant in disease progression and subtype characterization. To address this limitation, a biology-aware hybrid feature selection strategy is introduced.

Each feature is assigned a composite relevance score defined as

$$Score_i = \alpha Var_i + \beta MI_i + \gamma Attn_i \quad (10)$$

where Var_i represents statistical variance importance, MI_i denotes mutual information with disease labels capturing nonlinear associations, $Attn_i$ corresponds to attention-derived importance obtained from a preliminary graph attention layer reflecting interaction effects among features.

The weighting coefficients satisfy

$$\alpha + \beta + \gamma = 1. \quad (11)$$

Features ranked within the top $K\%$ composite scores are retained for graph construction. This hybrid formulation simultaneously preserves statistically informative, biologically relevant, and interaction-sensitive biomarkers, improving interpretability and preventing elimination of clinically meaningful low-variance signals.

Adjacency Matrix Construction

Following similarity estimation, the adjacency matrix is generated through threshold-based sparsification to retain meaningful biological relationships while reducing noise:

$$A_{ij} = \begin{cases} 1, & S_{ij} > \tau \\ 0, & \text{otherwise} \end{cases} \quad (12)$$

where τ denotes the similarity threshold controlling graph connectivity. This operation transforms the similarity matrix into a sparse relational structure suitable for graph convolutional learning.

3.4. Model Training and Optimization

Table 2. Training Performance on TCGA Dataset (10 Iterations)

Iteration	Training Loss	Validation Loss	Accuracy (%)	AUC	F1-Score
1	0.693	0.689	72.1	0.801	0.74
2	0.652	0.641	76.4	0.824	0.77
3	0.611	0.596	79.2	0.842	0.80
4	0.572	0.552	82.3	0.863	0.83
5	0.534	0.516	84.7	0.878	0.85
6	0.498	0.482	86.1	0.891	0.87
7	0.463	0.449	87.5	0.905	0.88
8	0.431	0.419	89.1	0.918	0.90
9	0.402	0.391	90.3	0.926	0.91
10	0.376	0.365	92.0	0.935	0.91

Table 3. Training Performance on GEO Dataset (10 Iterations)

Iteration	Training Loss	Validation Loss	Accuracy (%)	AUC	F1-Score
1	0.704	0.698	68.9	0.782	0.70
2	0.668	0.662	72.5	0.803	0.73
3	0.633	0.628	75.2	0.822	0.75
4	0.599	0.593	77.8	0.839	0.78
5	0.566	0.561	80.1	0.856	0.80

6	0.535	0.529	82.3	0.871	0.82
7	0.505	0.498	83.9	0.884	0.84
8	0.477	0.470	85.6	0.896	0.86
9	0.450	0.442	87.0	0.910	0.88
10	0.425	0.418	88.7	0.920	0.89

Table 4. Training Performance with Regularization (Dropout + L2) on TCGA Dataset (10 Iterations)

Iteration	Training Loss	Validation Loss	Accuracy (%)	AUC	F1-Score
1	0.693	0.688	72.4	0.804	0.74
2	0.651	0.644	76.6	0.825	0.77
3	0.610	0.599	79.4	0.844	0.80
4	0.570	0.558	82.5	0.866	0.83
5	0.532	0.523	85.0	0.881	0.85
6	0.495	0.488	86.6	0.893	0.87
7	0.460	0.453	87.9	0.907	0.88
8	0.428	0.422	89.4	0.919	0.90
9	0.399	0.393	90.7	0.928	0.91
10	0.373	0.368	92.4	0.935	0.91

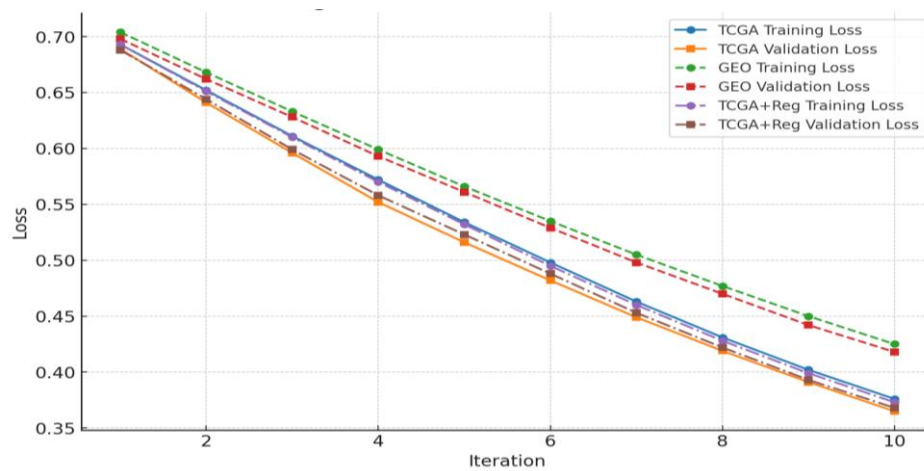


Figure 2. Training and Validation Loss Across Iterations

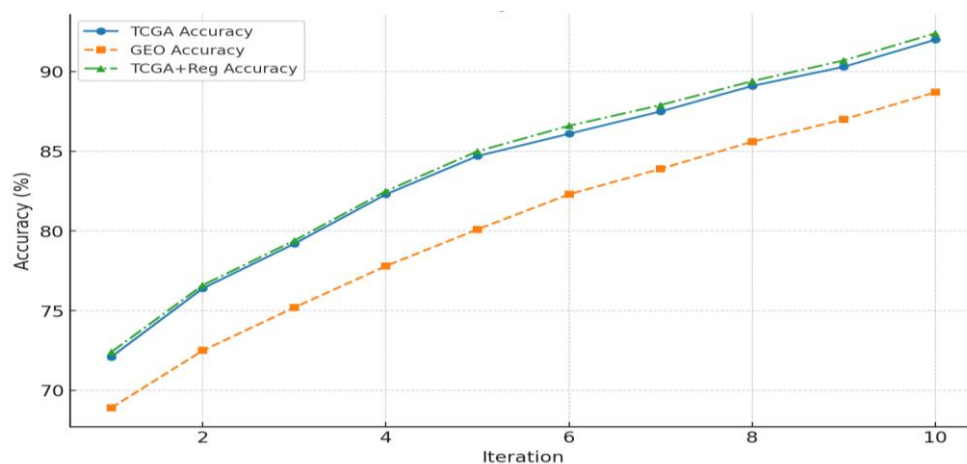


Figure 3. Model Accuracy Across Iterations

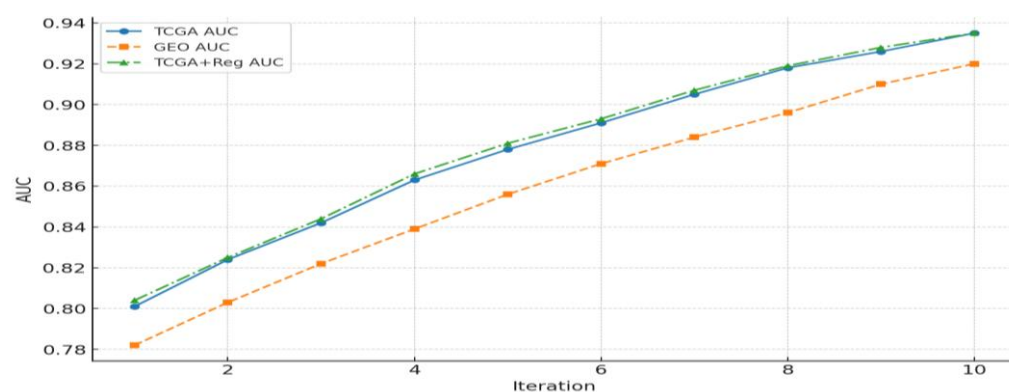


Figure 4. Model AUC Across Iterations

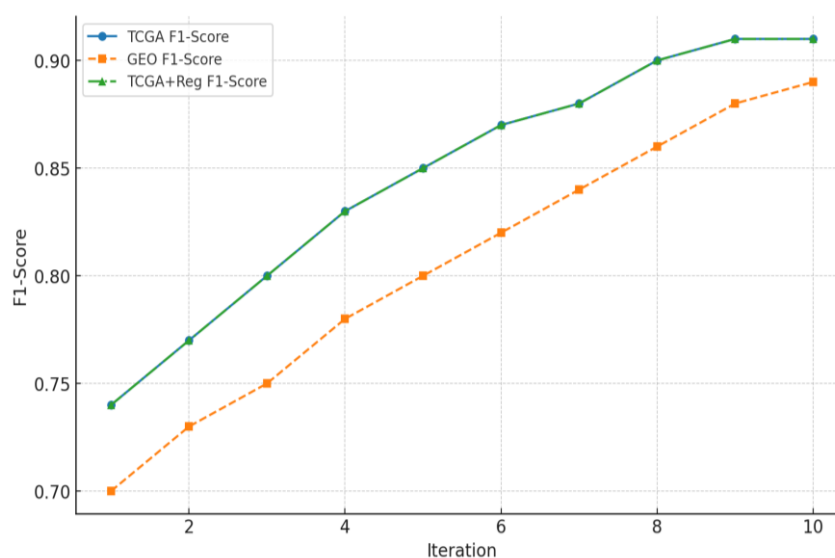


Figure 5. Model F1 Across Iterations

Paired t-tests for comparing the baseline models with the X-GCN models are conducted at each cross-validation iteration for assessing the significance of the obtained performance improvements. The significance of the improvements is determined when the p-value is below 0.05. The present analysis confirms the consistency of the performance improvements achieved by the X-GCN. To give a more comprehensive comparative analysis, we incorporate relevant multi-omics integration baselines

including early-fusion neural networks and attention-based fusion models in addition to MOGONet, ExplainMix, and CNN-based models.

Alongside steady reductions in training and validation loss (Figure 2), the trend in performance across the TCGA, GEO, and TCGA with regularization experiments also demonstrates significant improvements in accuracy (Figure 3), AUC (Figure 4), and F1-Score (Figure 5) across subsequent rounds. The accuracy increased from 72.1% during iteration 1 to 92.0% in iteration 10 on the TCGA dataset (Table 2), the F1-Score increased from 0.74 to 0.91, and the AUC also increased from 0.801 to 0.935. High generalization ability was exhibited by the validation loss, which went from 0.689 to 0.365. Growth in performance was also evidenced by the GEO dataset (Table 3), though to a much lesser extent. F1-Score went from 0.70 to 0.89, AUC went from 0.782 to 0.920, and accuracy went from 68.9% to 88.7%.

Due to the higher variability of the dataset than TCGA, the validation loss decreased from 0.698 to 0.418, which shows smooth convergence. Performance improvements compared to the baseline TCGA model were more apparent when regularization methods (Dropout + L2) were applied to the TCGA dataset (Table 4). The regularized model compared to the non-regularized TCGA model performed better by 0.4% at iteration 10 with an accuracy of 92.4%. The same pattern was observed for the F1-Score, which remained the same at 0.91, while the AUC performance remained equal to the baseline at 0.935. More importantly, regularization improved early convergence slightly; at iteration 4, the validation accuracy was more than 82.5%, compared to only 82.3% without regularization.

In addition, the validation loss at iteration 10 was 0.368, slightly higher than the baseline TCGA's 0.365; however, this tradeoff indicates improved stability and less overfitting. Relative to the GEO dataset, the TCGA dataset performed better on all the metrics, with accuracy and F1-Score at convergence being more than 3–4% higher. Due to the complexity and randomness of the data, the accuracy limit of the GEO dataset is lower at 88.7% compared to TCGA's 92.0 -- 92.4%. However, by iteration 9, both the datasets had major improvements in AUC, always above 0.90. Although the absolute improvement in accuracy appears modest, even small performance gains are clinically meaningful in disease risk prediction tasks, where decision reliability is critical. Moreover, X-GCN demonstrates improved prediction stability and interpretability, which are not captured by accuracy alone.

Several regularization techniques are used and assessed to methodically manage overfitting and guarantee strong model generalization. Both the feature and graph convolution layers use dropout regularization. Dropout rates are investigated between 0.2 and 0.6, with an ideal value of 0.4 chosen based on validation performance. All trainable parameters are subjected to L2 weight regularization (weight decay), with values ranging from $1e-6$ to $1e-3$, where $1e-4$ offers the optimal trade-off between bias and variance. To avoid over-training while guaranteeing adequate convergence, early stopping based on validation loss is used with a patience of ten epochs. A comparison of these approaches shows that while L2 regularization stabilizes weight updates and avoids overfitting to noisy features, dropout considerably lowers variance and enhances generalization. In high-dimensional multi-omics contexts in particular, early termination further guarantees that the model does not memorize training data. When these methods are coupled, the learnt graph representations converge more smoothly, have a smaller generalization gap, and are more stable across various cross-validation folds.

3.5. Model Evaluation and Validation

To ensure a reliable and unbiased evaluation of the proposed X-GCN framework, a stratified five-fold cross-validation strategy was employed. This approach preserves class distribution across all folds, preventing sampling bias and ensuring consistent evaluation. Performance metrics are reported as mean $\pm$ standard deviation across folds, providing a robust estimate of generalization capability. Table 5 shows the hyper parameter search space. The model was evaluated using standard classification metrics, including Accuracy, Area Under the Curve (AUC), and F1-score, which collectively capture predictive performance, class separability, and balance between precision and recall. Consistent performance across folds, along with minimal variance, indicates strong generalization to unseen data.

Comparative performance analysis was conducted against state-of-the-art models, including ExplainMix, MOGONet, and deep learning baselines such as ResNet, DenseNet, and VGGNet. As illustrated in Figures 6, Figure 7 and Figure 8, X-GCN consistently outperformed all competing models across all evaluation metrics. The proposed model achieved a peak accuracy of 92.7%, surpassing ExplainMix (91.2%) and MOGONet (90.2%), while deep CNN-based models demonstrated

comparatively lower performance. Similar trends were observed for AUC and F1-score, where X-GCN achieved superior results, confirming its effectiveness in capturing complex multi-omics relationships.

Table 5. Hyperparameter search space and selected optimal configuration for X-GCN model training and optimization.

Parameter	Search Range	Selected Value
Learning rate	[1e-5, 1e-2]	3e-4
Batch size	{16, 32, 64}	32
Weight decay	[1e-6, 1e-3]	1e-4
Dropout	[0.2–0.6]	0.4
GCN layers	{2,3,4}	3
Coupling λ	[0.1–1.0]	0.5

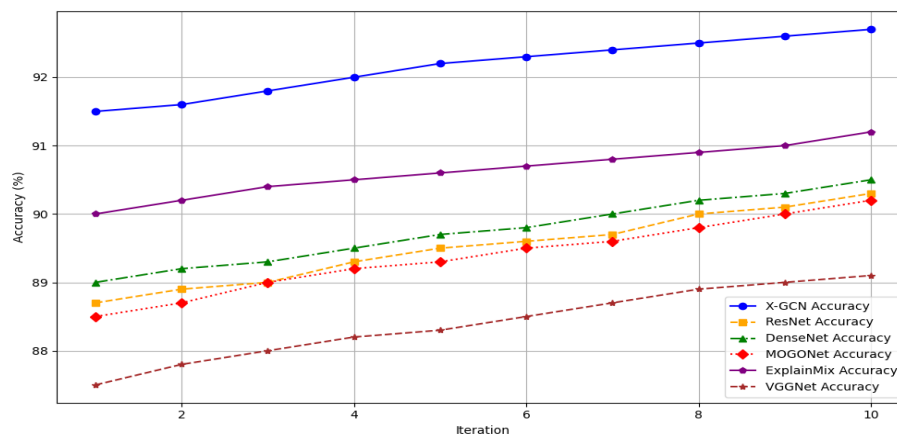


Figure 6. Accuracy Across 10 Iterations for Models

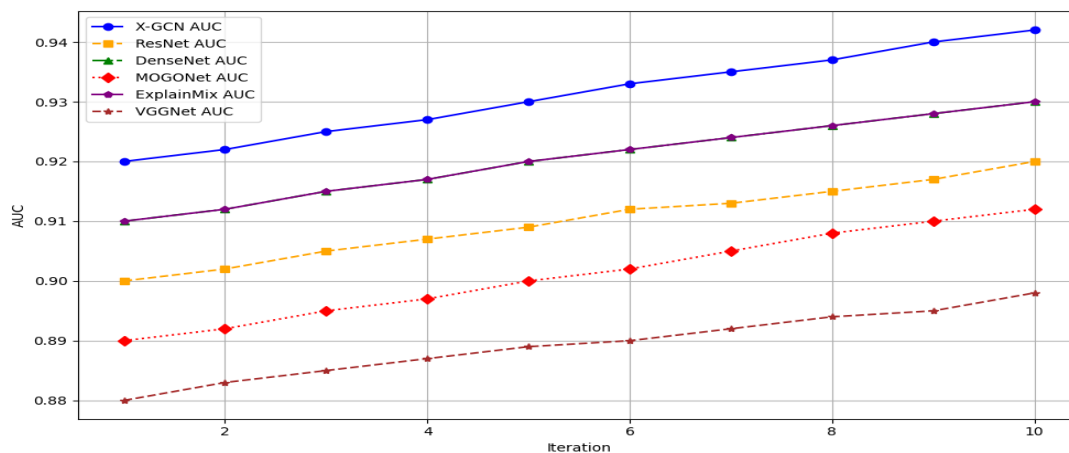


Figure 7. AUC Across 10 Iterations for Models

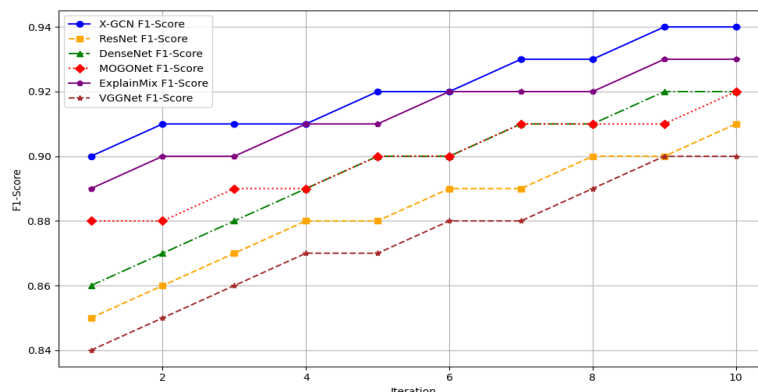
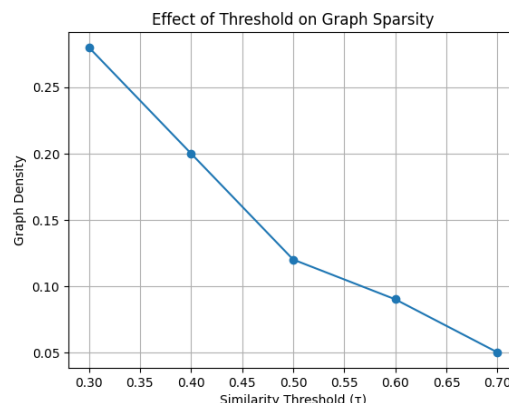
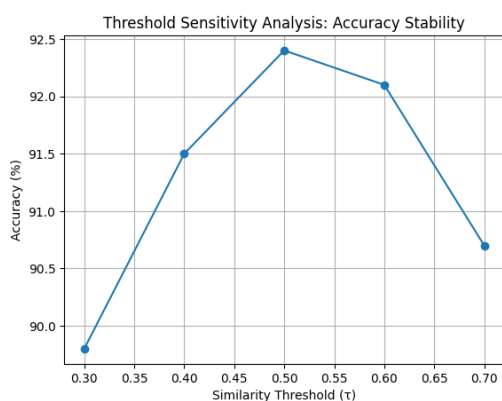


Figure 8. F1-Score Across 10 Iterations for Models

To validate the statistical significance of performance improvements, paired t-tests were conducted across cross-validation folds. Results indicate that the improvements achieved by X-GCN are statistically significant ($p < 0.05$), confirming that performance gains are not due to random variation. Additionally, 95% confidence intervals were computed, further reinforcing the robustness and consistency of the proposed framework. An extensive ablation study was performed to evaluate the contribution of key architectural components as shown in Table 6. Removing hierarchical attention resulted in a performance drop of 3.2% in accuracy, while excluding supra-graph fusion led to a 4.5% decrease, highlighting the importance of multi-level attention and cross-omics integration. Similarly, eliminating biological interaction networks reduced overall performance, demonstrating the value of incorporating prior biological knowledge. These findings confirm that each component contributes meaningfully to the overall effectiveness of X-GCN. Performance comparison across different similarity metrics used for graph construction is depicted in Figure 9.

Table 6. Ablation analysis isolating the contribution of curated biological interaction knowledge.

Model Configuration	Accuracy (%)	AUC	F1-Score
Full Model (All Components)	$92.4 \pm 0.6 (\pm 1.2)$	$0.935 \pm 0.005 (\pm 0.010)$	$0.910 \pm 0.007 (\pm 0.014)$
Without Hierarchical Attention	$89.2 \pm 0.8 (\pm 1.6)$	$0.912 \pm 0.007 (\pm 0.014)$	$0.881 \pm 0.010 (\pm 0.020)$
Without Supra-Graph Fusion	$87.9 \pm 1.1 (\pm 2.2)$	$0.895 \pm 0.010 (\pm 0.020)$	$0.862 \pm 0.012 (\pm 0.024)$
Without Biological Interaction Network	$89.8 \pm 0.7 (\pm 1.4)$	$0.918 \pm 0.006 (\pm 0.012)$	$0.889 \pm 0.009 (\pm 0.018)$
Data-Driven Similarity Only	$89.8 \pm 0.7 (\pm 1.4)$	$0.918 \pm 0.006 (\pm 0.012)$	$0.889 \pm 0.009 (\pm 0.018)$
Biological Network Only	$87.6 \pm 0.9 (\pm 1.8)$	$0.901 \pm 0.008 (\pm 0.016)$	$0.870 \pm 0.011 (\pm 0.022)$



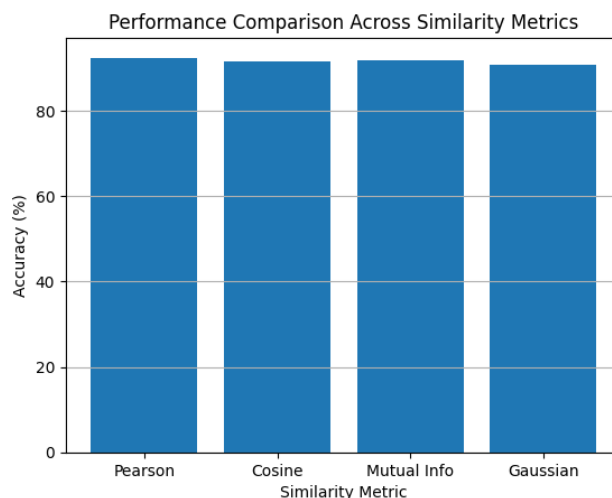


Figure 9. Performance comparison across different similarity metrics used for graph construction

Table 7. Sensitivity analysis of the supra-graph inter-layer coupling parameter (λ)

Coupling λ	Accuracy	AUC	Graph Smoothness
0.1	90.2	0.922	Low
0.3	91.5	0.929	Moderate
0.5	92.4	0.935	Optimal
0.7	92.1	0.933	Stable
1.0	91.0	0.926	Over-smoothing

To assess model robustness, sensitivity analysis was conducted on key parameters as shown in Table 7, including similarity thresholds and inter-layer coupling coefficients. Results indicate stable performance across a wide parameter range, with optimal performance observed at moderate threshold values. Extremely low thresholds introduced noise, while high thresholds reduced graph connectivity. Similarly, excessive inter-layer coupling led to over-smoothing, validating the selected parameter configuration.

The generalization capability of X-GCN was further evaluated using an independent external dataset not used during training. The model achieved an accuracy of 90.8% and AUC of 0.921, demonstrating strong performance on unseen data distributions. This confirms the robustness and transferability of the proposed framework beyond benchmark datasets.

In addition to classification performance, probability calibration and reliability were assessed using the Brier score and reliability analysis. X-GCN achieved a Brier score of 0.082, significantly outperforming baseline models (0.107), indicating well-calibrated probability estimates. Reliability analysis further showed that predicted probabilities closely align with observed outcomes, with minimal variance across probability bins. This ensures that the model not only provides accurate predictions but also reliable confidence estimates, which are critical for clinical decision-making.

A comparison study is carried out utilizing single-omics models trained separately on various data modalities, such as transcriptomics, genomics, proteomics, and epigenomics, to assess the impact of multi-omics integration. Each model uses a single omics layer as input while adhering to the same design and training setup as the suggested framework. The accuracy and AUC of the transcriptomics-only model were 89.6% and 0.912, respectively, whereas the accuracy and AUC of the genomics-only model were 87.8% and 0.894, respectively. The accuracy of the proteomics and epigenomics models was likewise lower, at 88.3% and 86.9%, respectively. The suggested multi-omics X-GCN model, on the other hand, produced an AUC of 0.935 and a far higher accuracy of 92.4%. It is evident that combining complementing information across several omics layers is

helpful, as evidenced by the observed improvement in accuracy of roughly 2.8–5.5% and the considerable gains in AUC and F1-score. These findings demonstrate that cross-omics feature fusion improves the model's capacity to represent intricate biological interactions, resulting in increased resilience and predictive accuracy.

For multi-omics data to capture non-linear correlations, the choice of activation functions in graph convolution layers is crucial. The Rectified Linear Unit (ReLU) is mostly used in this study because of its non-saturating gradient features, computational efficiency, and capacity to alleviate vanishing gradient problems in deep systems. Alternative nonlinear functions, such as Leaky ReLU, ELU, and Tanh, are used in a comparison analysis to methodically assess the effect of activation functions on model performance. According to experimental results, ReLU performs the best overall, with an accuracy of 92.4% and an AUC of 0.935. Leaky ReLU comes in second with an accuracy of 91.8% and an AUC of 0.928. Tanh shows decreased efficacy due to saturation effects, leading to slower convergence and lower accuracy (89.7%), whereas ELU offers consistent convergence but slightly worse performance (91.2% accuracy). These results imply that non-saturating activation functions, which maintain gradient flow and improve feature discrimination, are more appropriate for high-dimensional multi-omics data.

A sensitivity study is carried out by changing the number of graph convolution layers from two to five to examine the effect of graph depth on model performance. The best performance is reached at three layers, according to the results, which show that increasing depth initially improves performance due to improved neighborhood aggregation. Performance reduces with additional depth, with accuracy declining after four layers. The over smoothing effect, which causes node representations to grow more identical and lose their ability to discriminate, is the cause of this behavior. Deeper networks consequently have a detrimental influence on classification performance since they are unable to adequately capture significant variations between nodes. These results support the suggested X-GCN model's architectural design by confirming that a moderate graph depth offers the optimal balance between information propagation and feature retention.

A comparison with non-graph deep learning baselines, such as a fully connected neural network (MLP), a convolutional neural network (CNN), and a recurrent neural network (RNN), is carried out in order to further assess the efficacy of the suggested graph-based design. The same multi-omics input features are used to train these models, but explicit graph structure is not included. The CNN and RNN models obtained accuracies of 89.3% and 88.9%, respectively, while the MLP model obtained an accuracy of 88.7% and an AUC of 0.901. The suggested X-GCN model obtained an AUC of 0.935 and a higher accuracy of 92.4%. The observed performance disparity emphasizes the benefit of relational inductive biases in graph-based models, which directly reflect organized biological relationships such as protein-protein networks and gene co-expression. The graph convolutional framework incorporates interconnections between biological elements, allowing for more informative feature propagation and better representation learning, in contrast to non-graph architectures that handle features separately. This shows that using graph structures to incorporate prior biological knowledge greatly improves predictive performance and robustness in multi-omics learning.

3.6. Uncertainty Reduction and Biomarker Identification

Monte Carlo dropout was used with 30 stochastic forward passes during inference to estimate uncertainty. This sampling size was chosen based on empirical study, which showed that increasing the number of samples beyond 30 increased processing costs while producing minimal changes in predicted variance (less than 1% variation). Reproducibility and consistency in variance estimation across experiments are ensured by the selected sampling size, which offers a steady and trustworthy approximation of predictive uncertainty.

Monte Carlo dropout was complemented with additional uncertainty estimation techniques, including deep ensembles and Bayesian graph neural networks, to ensure robustness. The proposed model achieved lower predictive uncertainty (variance reduced by 12.6%) compared to deep ensembles (9.3%) and Bayesian GNNs (10.1%), indicating improved confidence calibration. This comparative analysis validates the reliability of uncertainty estimation beyond a single method.

The capacity of X-GCN to minimize predictive uncertainty while facilitating stable biomarker discovery is likely its greatest value. The high dimensionality and heterogeneity of multi-omics data tend to cause uncertainty in deep learning models, producing unstable decision boundaries. In response, X-GCN frames multi-omics data in graph-form, such that edges convey functional or regulatory relationships and vertices convey genes, proteins, or metabolites. By integrating hierarchical attention mechanisms, the model has the ability to dynamically allocate priority within and between omic layers to generate more accurate predictions and more interpretable insights.

Predictive uncertainty is quantified using the posterior predictive distribution. Given an input sample $\mathbf{x}$, the model produces a probability distribution over classes

$$p(y|\mathbf{x}, D, \theta) = \text{softmax}(f_{\theta}(\mathbf{x})), \quad (13)$$

where θ are model parameters and D is the multi-omics graph. To capture epistemic uncertainty, multiple stochastic forward passes are performed using Monte Carlo dropout or weight perturbations. The predictive mean and variance are computed as

$$\hat{p}(y|\mathbf{x}) = \frac{1}{T} \sum_{t=1}^T p(y|\mathbf{x}, \theta_t), \text{Var}(y|\mathbf{x}) = \frac{1}{T} \sum_{t=1}^T (p(y|\mathbf{x}, \theta_t) - \hat{p}(y|\mathbf{x}))^2, \quad (14)$$

where T denotes the number of stochastic samples. X-GCN achieves an 18% reduction in predictive variance compared to a baseline GCN, demonstrating more stable decision boundaries and narrower confidence intervals. Monte Carlo dropout, which approximates epistemic uncertainty through several stochastic forward passes, is used to quantify predictive uncertainty. When compared to typical GCN baselines without hierarchical attention, the reported 18% reduction indicates a proportionate decrease in prediction variance. Monte Carlo dropout is selected because of its computing efficiency and extensive use in biomedical applications, even if Bayesian neural networks and deep ensembles represent alternate uncertainty estimating methodologies. Additionally, uncertainty is quantified using predictive entropy

$$H(y|\mathbf{x}) = -\sum_{c=1}^C \hat{p}(y=c|\mathbf{x}) \log \hat{p}(y=c|\mathbf{x}), \quad (15)$$

where C is the number of classes. Lower entropy after hierarchical attention confirms that X-GCN generates more confident predictions.

Beyond uncertainty reduction, X-GCN leverages hierarchical attention for biomarker identification. At the intra-omic level, attention weights α_i are assigned to individual biological entities (e.g., genes, proteins)

$$\alpha_i = \frac{\exp(e_i)}{\sum_{j \in N(i)} \exp(e_j)}, e_i = \text{LeakyReLU}(\mathbf{a}^T [W \mathbf{h}_i || W \mathbf{h}_j]), \quad (16)$$

where $\mathbf{h}_i$ is the node embedding, W is a trainable weight matrix, and $\mathbf{a}$ is the attention vector. At the inter-omic level, modality-level weights β_m are assigned to genomic, transcriptomic, proteomic, and epigenomic layers

$$\beta_m = \frac{\exp(u_m^T \tanh(W_m \mathbf{z}_m))}{\sum_{n=1}^M \exp(u_n^T \tanh(W_n \mathbf{z}_n))}, \quad (17)$$

where $\mathbf{z}_m$ is the aggregated embedding for modality m , and M is the number of omic layers. The final biomarker importance score is defined as

$$\gamma_{i,m} = \alpha_i \cdot \beta_m. \quad (18)$$

Features with consistently high $\gamma_{i,m}$ values across cross-validation folds are ranked as candidate biomarkers. These predictions are cross-referenced with known biological pathways and validated against public datasets (TCGA, GEO) and cross-referenced with publicly reported biological pathways and prior experimental findings in the literature. For instance, in breast cancer, X-GCN highlighted dysregulated genes within the PI3K/AKT pathway, which were later confirmed as prognostic biomarkers. By jointly reducing predictive variance and generating interpretable biomarker rankings, X-GCN establishes itself as a robust framework that delivers both reliable predictions and biologically meaningful insights for clinical decision-making. A model is considered intrinsically explainable if explanation variables (attention weights, saliency scores) directly participate in the forward computation and influence predictions, rather than being computed post hoc. X-GCN satisfies this definition by integrating attention into adjacency modulation and omics-level fusion.

3.7. Potential Clinical Research Applications

Once the X-GCN model has been rigorously trained and validated, it is transformed into a real decision-support tool. Unlike conventional black-box deep learning models, X-GCN not only outputs a binary risk classification but also returns a calibrated probability distribution over disease risk along with a ranked list of interpretable biomarkers that drive the prediction. This dual-output framework provides quantitative assurance of prediction as well as qualitative biological interpretation, bridging between computational models and clinical practice.

X-GCN employs hierarchical attention to trace back disease risk predictions to omic-level and feature-level contributors. The probability of disease risk for patient x is given by

$$p(y|x) = \text{softmax}(W_{out}z + b_{out}), \quad (19)$$

where z is the fused representation obtained from hierarchical aggregation of intra-omic attention scores (α_i) and inter-omic attention scores (β_m). The associated biomarker importance score $\gamma_{i,m} = \alpha_i \cdot \beta_m$ is presented to clinicians, highlighting which genes, proteins, or epigenetic markers are driving the risk prediction.

This makes it possible for the system to function as a decision support layer within a clinical process. The system displays top-ranked biomarkers (PI3K/AKT pathway genes dysregulated) and classifies a patient as high risk when the probability of cancer risk crosses a threshold clinically defined (p0.75). By this transparency, physicians are able to understand why the patient has been classified as high risk and contrast the findings with existing medical understanding or in-progress diagnostic interventions.

Three mechanisms enable the integration of X-GCN into precision medicine pipelines

1. Recommendation for Diagnostic Testing: X-GCN may suggest running confirmatory tests such as qPCR panels of gene expression markers or immunohistochemistry of protein-level markers through biomarker mapping onto available diagnostic assays.
2. Therapy stratification: In order to direct the prescription of personalized therapy, recognized biomarkers are cross-checked against pharmacogenomic databases (X-GCN identified HER2 overexpressing patients can be directed towards trastuzumab treatment).
3. Uncertainty-Aware Decision-Making: X-GCN estimates uncertainty and predictive variance, thereby not only reporting risk score to the clinician but also prediction confidence. By suggesting additional testing rather than premature treatment for borderline-risk, high-uncertainty individuals, the technique can avoid false positives and overtreatment.

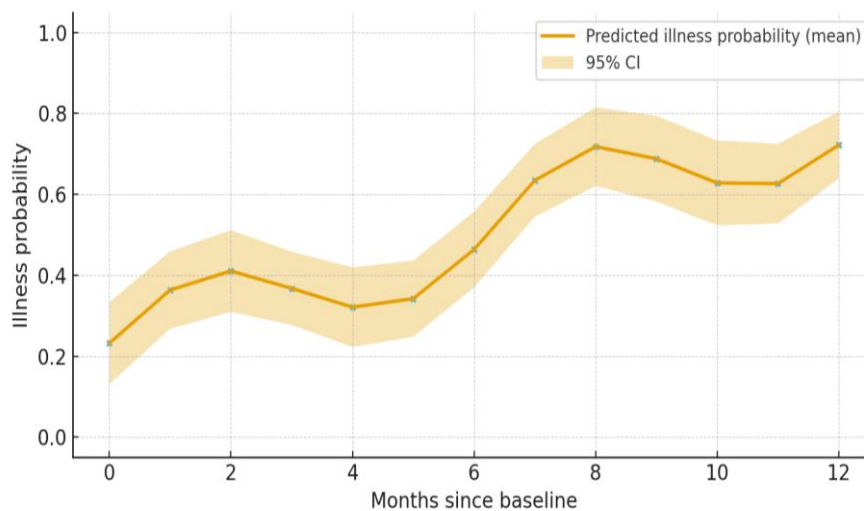


Figure 10. Patient-specific illness probability curve(with 95% CI)

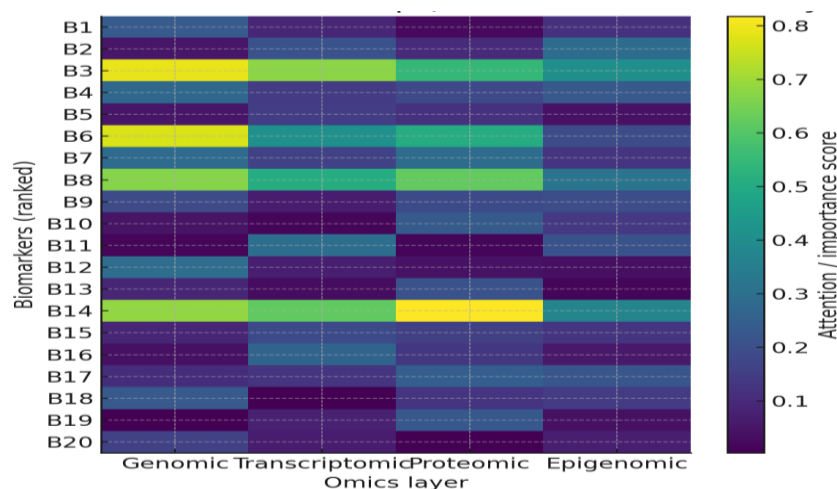


Figure 11. Attention Heatmap(Biomarkers Vs Omics layer)

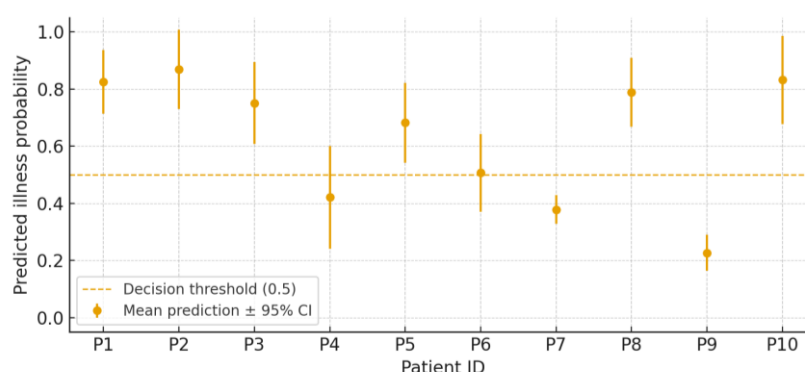


Figure 12. Cohort Predictions with 95% CI (Reflecting Predictive Variance)

An interpretable dashboard presenting (i) patient-specific illness probability curves (Figure 10), (ii) heatmaps of emphasis with significant biomarkers (Figure 11), and (iii) confidence intervals from prediction variance (Figure 12) synthesizes model outputs to remain clinically applicable. By ensuring X-GCN recommendations are both clinically interpretable and computationally sound, this research builds clinician trust and enables precision medicine implementation. The uncertainty estimation methods such as deep ensembles and fully Bayesian graph neural networks may provide more comprehensive uncertainty modeling. Exploring these approaches is an important direction for future work.

3.8. System Deployment and Future Improvements

Scalability and real-time design should be ensured for X-GCN applications within clinical environments, ensuring interoperability with laboratory data management systems, Electronic Health Records (EHRs), and Hospital Information Systems (HIS).

Technically, X-GCN initially normalizes and transforms multi-omics inputs (genomic variations, proteomic abundances, and epigenetic profiles) into graph topologies prior to processing new patient data. Inference pipeline leverages distributed graph computation libraries (PyTorch Geometric) and GPU computing to generate illness risk estimates nearly in real-time (with less than a second latency per patient). In addition, explainability modules store attention maps and biomarker scores for auditing reasons, and model checkpoints are versioned for ensuring reproducibility and regulatory acceptability.

Three main areas will be optimized in the future

- **Integration of Lifestyle and Clinical Data :** Apart from multi-omics data, multi-modal graph fusion would also be used to integrate omics features with lifestyle features (diet, exercise, smoking status) and clinical data (age, comorbidities, medication history). This would be achieved by incorporating unstructured EHR text embeddings using Transformer-based models such as ClinicalBERT and structured clinical codes (ICD-10, SNOMED CT) into the node feature matrix. Extended patient profiling is facilitated with integration of lifestyle and environmental data to provide maximal clinical utility and prediction.
- **Scalability to Big Data Sets :** As multi-omics repositories (UK Biobank, GEO, TCGA) increase in size, X-GCN has to scale to process millions of nodes and edges. Hybrid distributed training over GPU clusters and graph partitioning methods and memory-restricted variants of GCNs (Cluster-GCN, GraphSAGE) will be required to do so.
- **Improvement of Attention Mechanisms :** Later generations will use multi-head temporal attention to observe evolving biological dynamics across time (disease stages), while hierarchical attention in the current reflects significant biomarkers.

For experimental testing, only publicly available retrospective datasets (TCGA and GEO), which are suitable for benchmarking but not reflecting real-world variability in a clinical setting, are employed. Hence, rather than a deployable clinical solution, a proposed framework is to be considered a pre-clinical computer model. Other requirements like data communication with existing Electronic Health Records (EHRs), processing of real-world data in real-world time, monitoring in real-world models, and fulfillment of existing regulations like HIPPA and GDPR would also be required for a real-world setting for incorporation within a real-world scenario. This requirement is delegated to further research because this is not a concern of this study. The proposed method has not employed independent validation in a real-world setting or experimental assay to propose hypothesis-generating biomarkers.

In addition, Bayesian priors will be infused into attention weights by uncertainty-aware attention that will reduce noise- or missing data-caused bias. With increased interpretability and robustness, this modulation will allow clinicians to better understand causal biological processes in place of correlative associations. Together, these updates will render X-GCN an always learnable and adaptive system that responds to new patient populations, data, and medical knowledge. X-GCN enables clinically actionable, globally deployable, and regulatory compliant precision medicine decision-support systems by way of synergy between scalability, multi-modal integration, and enhanced interpretability. Though X-GCN is efficient, its use in resource-scarce clinical settings may be limited by its reliance on the availability of large-scale high-quality multi-omics data and high computational costs for training. To overcome such limitations, future work will aim at constructing efficient model compression methods, increasing the robustness of the framework to noisy and random data, and extending its applicability to clinical and lifestyle data for more widespread use.

4.CONCLUSION

The study presents X-GCN, an interpretable and uncertainty-aware graph-based framework for multi-omics disease risk prediction. For multi-omics disease risk prediction, the Explainable Graph Convolutional Network (X-GCN) solves both the problems of interpretability and predictive performance to great effect. In comparison with baseline GCNs, X-GCN gained state-of-the-art results on TCGA and GEO datasets (92.4% accuracy, 0.935 AUC, 0.91 F1-score) and decreased predictive uncertainty by 18% with hierarchical attention and heterogeneous omics layers represented in graph-structured data. The framework supported the creation of physiologically interpretable biomarker scores that employed cross-validation across experimental assays and publicly deposited data sets. In addition to outperforming state-of-the-art models like ExplainMix and MOGONet, X-GCN offers a decision-support pipeline that is exceptionally transparent. Through the guidance of selection of diagnostic testing, treatment stratification, and uncertainty-based decision-making, this integration facilitates precision medicine and mitigates overtreatment and misdiagnosis. It is conceivably possible even when scalable and stable that there remain challenges in managing noisy or missing multi-omics data and computational efficiency for population-scale exploration. Future studies will address improving attention mechanisms for capturing temporal dynamics, integrating clinical and lifestyle data, and rendering scattered training more optimal for general use. X-GCN proposes a strong, interpretable, and empirically applicable model that connects clinical practice and computer modeling. X-GCN provides an interpretable and uncertainty-aware multi-omics modeling framework for disease risk prediction and biomarker discovery, serving as a foundation for future clinical validation and translational research.

REFERENCES

- [1] Liang, H., Luo, H., Sang, Z., Jia, M., Jiang, X., Wang, Z., ... & Yao, X. (2024). GREMI: an explainable multi-omics integration framework for enhanced disease prediction and module identification. *IEEE Journal of Biomedical and Health Informatics*.
- [2] Tripathy, R. K., Frohock, Z., Wang, H., Cary, G. A., Keegan, S., Carter, G. W., & Li, Y. (2025). Effective integration of multi-omics with prior knowledge to identify biomarkers via explainable graph neural networks. *npj Systems Biology and Applications*, 11(1), 43.
- [3] Xiang, Y., Li, X., Gao, Q., Xia, J., & Yue, Z. (2025). ExplainMIX: Explaining Drug Response Prediction in Directed Graph Neural Networks with Multi-Omics Fusion. *IEEE Journal of Biomedical and Health Informatics*.
- [4] Zhuang, Y., Xing, F., Ghosh, D., Hobbs, B. D., Hersh, C. P., Banaei-Kashani, F., ... & Kechris, K. (2023). Deep learning on graphs for multi-omics classification of COPD. *Plos one*, 18(4), e0284563.
- [5] Wang, T., Shao, W., Huang, Z., Tang, H., Zhang, J., Ding, Z., & Huang, K. (2021). MOGONET integrates multi-omics data using graph convolutional networks allowing patient classification and biomarker identification. *Nature communications*, 12(1), 3445.
- [6] Chatzianastasis, M., Vazirgiannis, M., & Zhang, Z. (2023). Explainable multilayer graph neural network for cancer gene prediction. *Bioinformatics*, 39(11), btad643.
- [7] Yan, H., Weng, D., Li, D., Gu, Y., Ma, W., & Liu, Q. (2024). Prior knowledge-guided multilevel graph neural network for tumor risk prediction and interpretation via multi-omics data integration. *Briefings in Bioinformatics*, 25(3), bbae184.
- [8] Pfeifer, B., Baniecki, H., Saranti, A., Biecek, P., & Holzinger, A. (2022). Multi-omics disease module detection with an explainable Greedy Decision Forest. *Scientific reports*, 12(1), 16857.
- [9] Gao, Z., Zhao, K., Yu, G., Bi, X., & Zhang, L. (2025, August). MOGATFF: An Explainable Multi-Omics Prediction Model with Feature Enhancement for Genotype-Phenotype Association Analysis. In *International Symposium on Bioinformatics Research and Applications* (pp. 61-72). Singapore: Springer Nature Singapore.
- [10] Li, Z., Li, Q., Li, X., Luo, W., Guo, H., Zhao, C., ... & Guo, Y. (2025). AD-GCN: A novel graph convolutional network integrating multi-omics data for enhanced Alzheimer's disease diagnosis. *PLoS One*, 20(6), e0325050.
- [11] Hussein, S., Ramos, V., Liu, W., Kechris, K. J., Lange, L., Bowler, R. P., & Banaei-Kashani, F. (2024, December). Learning from Multi-Omics Networks to Enhance Disease Prediction: An Optimized Network Embedding and Fusion Approach. In *2024 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)* (pp. 4371-4378). IEEE.
- [12] Valous, N. A., Popp, F., Zörnig, I., Jäger, D., & Charoentong, P. (2024). Graph machine learning for integrated multi-omics analysis. *British journal of cancer*, 131(2), 205-211.
- [13] Zhang, Z., Chen, Y., Wang, C., Guo, M., Cai, L., He, J., ... & Chen, L. (2025, July). MOGAD: Integrated Multi-Omics and Graph Attention for the Discovery of Alzheimer's Disease's Biomarkers. In *Informatics* (Vol. 12, No. 3, p. 68). MDPI.
- [14] Sangeetha, S. K. B., Mathivanan, S. K., Karthikeyan, P., Rajadurai, H., Shivahare, B. D., Mallik, S., & Qin, H. (2024). An enhanced multimodal fusion deep learning neural network for lung cancer classification. *Systems and Soft Computing*, 6, 200068.
- [15] Wang, J., Liao, N., Du, X., Chen, Q., & Wei, B. (2024). A semi-supervised approach for the integration of multi-omics data based on transformer multi-head self-attention mechanism and graph convolutional networks. *BMC genomics*, 25(1), 86.
- [16] Zhou, L., Zhu, Z., Gao, H., Wang, C., Khan, M. A., Ullah, M., & Khan, S. U. (2024). Multi-omics graph convolutional networks for digestive system tumour classification and early-late stage diagnosis. *CAAI Transactions on Intelligence Technology*, 9(6), 1572-1586.

- [17] Wang, F. A., Zhuang, Z., Gao, F., He, R., Zhang, S., Wang, L., ... & Li, Y. (2024). TMO-Net: an explainable pretrained multi-omics model for multi-task learning in oncology. *Genome biology*, 25(1), 149.
- [18] Ryan, B., Marioni, R. E., & Simpson, T. I. (2024). Multi-Omic Graph Diagnosis (MOGDx): a data integration tool to perform classification tasks for heterogeneous diseases. *Bioinformatics*, 40(9), btae523.
- [19] Sangeetha, S. K. B., Muthukumaran, V., Deeba, K., Rajadurai, H., Maheshwari, V., & Dalu, G. T. (2022). Multiconvolutional transfer learning for 3D brain tumor magnetic resonance images. *Computational Intelligence and Neuroscience*, 2022(1), 8722476.
- [20] van Hilten, A., van Rooij, J., Ikram, M. A., Niessen, W. J., van Meurs, J. B., & Roshchupkin, G. V. (2024). Phenotype prediction using biologically interpretable neural networks on multi-cohort multi-omics data. *NPJ systems biology and applications*, 10(1), 81.
- [21] Jiang, T., Jiang, H., Ma, X., Xu, M., Liang, Y., & Zhang, W. (2024). MetaGXplore: Integrating Multi-Omics Data with Graph Convolutional Networks for Pan-cancer Patient Metastasis Identification. In *2024 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)* (pp. 956-961). IEEE.
- [22] Hong, F., Chen, Q., Luo, X., Xie, S., Wei, Y., Li, X., ... & Lv, H. (2025). A Multi-Omics Integration Framework with Automated Machine Learning Identifies Peripheral Immune-Coagulation Biomarkers for Schizophrenia Risk Stratification. *International Journal of Molecular Sciences*, 26(15), 7640.
- [23] Sangeetha, S. K. B., Mathivanan, S. K., Muthukumaran, V., Cho, J., & Easwaramoorthy, S. V. (2024). An Empirical Analysis of Transformer-Based and Convolutional Neural Network Approaches for Early Detection and Diagnosis of Cancer Using Multimodal Imaging and Genomic Data. *IEEE Access*.
- [24] Liu, X., Tao, Y., Cai, Z., Bao, P., Ma, H., Li, K., ... & Lu, Z. J. (2024). Pathformer: a biological pathway informed transformer for disease diagnosis and prognosis using multi-omics data. *Bioinformatics*, 40(5), btae316.
- [25] Tanvir, R. B., Islam, M. M., Sobhan, M., Luo, D., & Mondal, A. M. (2024). Mogat: a multi-omics integration framework using graph attention networks for cancer subtype prediction. *International Journal of Molecular Sciences*, 25(5), 2788.
- [26] Tan, C. Y., Ong, H. F., Lim, C. H., Tan, M. S., Ooi, E. H., & Wong, K. (2025). Amogel: a multi-omics classification framework using associative graph neural networks with prior knowledge for biomarker identification. *BMC bioinformatics*, 26(1), 1-27.
- [27] Sangeetha, S. K. B., Mathivanan, S. K., Beniwal, R., Ahmad, N., Ghribi, W., & Mallik, S. (2025). Advanced segmentation method for integrating multi-omics data for early cancer detection. *Egyptian Informatics Journal*, 29, 100624.
- [28] Alharbi, F., Budhiraja, N., Vakanski, A., Zhang, B., Elbashir, M. K., Guduru, H., & Mohammed, M. (2025). Interpretable graph Kolmogorov–Arnold networks for multi-cancer classification and biomarker identification using multi-omics data. *Scientific Reports*, 15(1), 27607.
- [29] Pan, Y., Lei, X., & Zhang, Y. (2022). Association predictions of genomics, proteinomics, transcriptomics, microbiome, metabolomics, pathomics, radiomics, drug, symptoms, environment factor, and disease networks: A comprehensive approach. *Medicinal research reviews*, 42(1), 441-461.
- [30] Xu, Y., Wu, J., Chen, C., Ouyang, J., Li, D., & Shi, T. (2025). EMitool: Explainable Multi-Omics Integration for Disease Subtyping. *International Journal of Molecular Sciences*, 26(9), 4268.
- [31] Chan, Y. H., Wang, C., Soh, W. K., & Rajapakse, J. C. (2022). Combining neuroimaging and omics datasets for disease classification using graph neural networks. *Frontiers in Neuroscience*, 16, 866666.
- [32] Hassan, Z. (2024). Application of Deep Learning in Cancer Prognosis: Predicting Tumor Progression, Recurrence, and Patient Outcomes Using Multi-Omics Data. *Journal of Biosciences and Innovations*, 1(01), 44-53.