



## AI-Powered Drug Repurposing using Graph Neural Networks (GNNs)

Vedant Ghodake<sup>1\*</sup>, Mahesh Bhandari<sup>2</sup>, Sayali Gaikwad<sup>3</sup>, Ishan Garud<sup>4</sup>, Sonali Ghule<sup>5</sup>, Rutika Harde<sup>6</sup>

<sup>1,2,3,4,5,6</sup> Department of Information Technology, Vishwakarma Institute of Technology, Pune, India

\*Correspondence to: Vedant Ghodake, Department of Information Technology, Vishwakarma Institute of Technology, Pune, India; E-mail: vedant.ghodake241@vit.edu

Received: July 03, 2026; Manuscript No: JADS-26-3794; Editor Assigned: July 07, 2026; PreQc No: JADS-26-3794 (PQ); Reviewed: July 13, 2026; Revised: July 21, 2026; Manuscript No: JADS-26-3794 (R); Published: August 26, 2026

Citation: Ghodake V, Bhandari M, Gaikwad S, Garud I, Ghule S, Harde R (2026). AI-Powered Drug Repurposing using Graph Neural Networks (GNNs). Adv. Data Sci. Comput. Intell. Vol.1 Iss.1, August (2026), pp:1-12.

Copyright: © 2026 Vedant Ghodake, Mahesh Bhandari, Sayali Gaikwad, Ishan Garud, Sonali Ghule, Rutika Harde. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

### ABSTRACT

Drug repurposing-the identification of novel therapeutic indications for existing, clinically approved compounds-has emerged as a cost-effective and time-efficient alternative to de novo drug discovery, particularly in the context of rare diseases, emerging pathogens, and treatment-resistant conditions. Conventional repurposing pipelines, however, remain constrained by the combinatorial vastness of the chemical-biological interaction space and the inherent limitations of similarity-based or single-omics inference methods. This paper introduces DeepCure Pro, a novel computational framework that leverages Graph Neural Networks (GNNs) to model and predict drug-disease associations within a heterogeneous biomedical knowledge graph integrating drug-target interactions, protein-protein interaction networks, gene-disease associations, and pathway-level annotations. DeepCure Pro employs a multi-relational graph attention architecture augmented with inductive node embedding strategies, enabling the framework to generalize to previously unseen drug and disease entities. The model is trained and validated using benchmark datasets, including DrugBank, DisGeNET, and STRING, with performance assessed via area under the receiver operating characteristic curve (AUROC), precision-recall metrics, and case-study validation against literature-confirmed repurposing events. Experimental results demonstrate that DeepCure Pro substantially outperforms traditional matrix factorization and network-propagation baselines, achieving superior predictive accuracy while maintaining biological interpretability through attention-weight analysis. This work contributes a scalable, extensible, and interpretable GNN-based architecture to the network medicine literature, offering significant implications for accelerating precision pharmacology and pandemic-response drug discovery pipelines.

**Keywords:** Graph Neural Networks; Drug Repurposing; GraphSAGE; Explainable Artificial Intelligence; Retrieval-Augmented Generation; Biomedical Knowledge Graphs; Link Prediction; Heterogeneous Graph Learning Intelligence

### INTRODUCTION

#### The Economics of Pharmaceutical Failure: Eroom's Law

The productivity of the pharmaceutical research and development (R&D) enterprise has, paradoxically, degraded in near-perfect inverse proportion to the exponential gains observed in computational hardware over the past five decades. This phenomenon, termed Eroom's Law - "Moore's Law" spelled backward -

describes the empirical observation that the number of novel therapeutics approved per billion dollars of inflation-adjusted R&D expenditure has halved approximately every nine years since 1950 [1]. Where semiconductor fabrication has benefited from compounding efficiency, drug discovery exhibits compounding cost inflation: a single approved molecular entity now demands a capitalized investment exceeding \$2.5 billion, once failure-adjusted opportunity costs and time-value-of-capital are incorporated into the total pipeline expenditure [2].

This economic burden is structurally rooted in the staged, high-attrition architecture of the discovery-to-approval pipeline. A candidate compound must sequentially clear target identification, lead optimization, preclinical toxicology, and three phases of clinical trials - a process typically spanning 10-15 years. At each transition, particularly the passage from Phase I safety trials into Phase II efficacy trials, attrition is severe; aggregate clinical success rates across all therapeutic areas remain below 10%, meaning that more than 90% of compounds entering human trials never reach market approval [3]. Because sunk costs from failed candidates are absorbed into the pricing and prioritization of surviving therapeutics, this compounding failure rate is the principal driver of Eroom's Law. Consequently, the traditional de novo discovery paradigm is structurally unsuited to addressing rare diseases, orphan indications, or rapidly emergent pathogens, where neither the time horizon nor the market size justifies the capital outlay.

Drug repurposing (also termed drug repositioning) - the systematic identification of new therapeutic indications for compounds already possessing established pharmacokinetic and safety profiles - circumvents the most expensive and failure-prone stages of the pipeline. Since repurposed candidates have typically already cleared Phase I toxicology, repurposing pipelines can, in principle, compress discovery-to-clinic timelines to 3-5 years at a fraction of de novo cost [4]. The central bottleneck is no longer chemical synthesis or toxicological safety, but rather hypothesis generation: the combinatorial search problem of identifying which of the thousands of approved compounds is mechanistically plausible against which of the thousands of characterized disease phenotypes.

### Network Medicine and the Link Prediction Formalism

The theoretical foundation for computational repurposing is provided by the discipline of Network Medicine, which posits that human disease is rarely attributable to an isolated gene defect, but instead emerges from perturbations propagating across an interconnected molecular interactome comprising protein-protein interactions, gene regulatory circuits, and metabolic pathways [5]. Under this paradigm, drugs, genes, proteins, and diseases are naturally represented not as isolated tabular records but as nodes within a heterogeneous biomedical graph  $G = (V, E, \phi, \psi)$ , where  $V$  is the set of entities,  $E$  the set of relational edges, and  $\phi: V \rightarrow A$ ,  $\psi: E \rightarrow R$  are type-mapping functions assigning each node and edge to a semantic class within node-type set  $A$  (e.g., drug, gene, disease) and relation-type set  $R$  (e.g., targets, treats, associated-with).

Within this formalism, the drug repurposing problem is reducible to a link prediction task: given the observed subgraph  $G_{obs}$  inside  $G$ , learn a scoring function  $f_{\theta}(u, v) \rightarrow [0,1]$  that estimates the likelihood of an unobserved therapeutic edge  $(u,v)$  belonging to  $E_{treats}$

existing between drug node  $u$  and disease node  $v$ . This reformulation transforms an intractable biochemical search space into a well-posed statistical inference problem, amenable to graph representation learning.

### Four Critical Limitations of Prior Computational Approaches

#### Data Heterogeneity

Biomedical knowledge graphs are inherently multi-relational and multi-typed: a single graph simultaneously encodes physically distinct interaction semantics - a drug-targets-protein edge is not physically or statistically equivalent to a disease-associated-with-gene edge. Homogeneous graph models, which collapse all edges into a single undifferentiated adjacency structure, discard this relational semantics, conflating mechanistically distinct pathways and inducing severe information loss during message aggregation [6].

#### Hub Bias - The Super-Node Problem

Biomedical interaction networks are empirically scale-free, following a power-law degree distribution in which a small subset of nodes - such as p53, TP53-associated pathways, or promiscuous kinase targets - accumulate disproportionately high connectivity [7]. During neighborhood aggregation, these "super-nodes" dominate the message-passing signal, causing representation learning to become biased toward generic hub-adjacency patterns rather than specific, mechanistically meaningful substructures - a phenomenon closely analogous to oversmoothing in deep GNN literature [8].

#### The Interpretability Gap

Deep graph architectures, by virtue of their distributed, non-linear parameterization across multiple aggregation layers, function as opaque predictive engines. In a clinical or pharmacological decision context, a bare probability score  $f_{\theta}(u,v) = 0.91$  carries negligible actionable value absent a mechanistic rationale; regulatory and scientific stakeholders require causal or associative justification traceable to established biological literature before a computational hypothesis can be escalated to costly wet-lab validation [9]. This "black-box" deficiency remains the principal adoption barrier for GNN-based repurposing tools in translational pipelines.

#### The Cold-Start Problem

Classical transductive graph embedding techniques (e.g., matrix factorization, DeepWalk, node2vec) compute a fixed embedding table indexed to nodes observed during training. Such methods are structurally incapable of generating representations for novel compounds or newly characterized targets introduced after training convergence, without complete model retraining [10]. Given that pharmaceutical databases are continuously updated with newly approved and experimental compounds, this

transductive constraint renders many prior systems operationally obsolete shortly after deployment.

### DeepCure Pro: A Structural-Semantic Hybrid Solution

To resolve these four limitations concurrently, this paper proposes DeepCure Pro, a hybrid computational framework that couples an inductive, heterogeneous message-passing backbone with a retrieval-grounded semantic reasoning layer. At the structural level, DeepCure Pro employs relation-aware GraphSAGE aggregation over the Stanford BioSNAP heterogeneous biomedical network, enabling generalizable embedding synthesis for previously unseen drug and gene entities through inductive neighborhood sampling - directly resolving the cold-start limitation. At the semantic level, the framework appends a Retrieval-Augmented Generation (RAG) explainability module, orchestrating Gemini 1.5 Pro against real-time literature retrieval via the Semantic Scholar API, converting raw link-prediction scores into citation-grounded mechanistic narratives. This dual-layer architecture - structural inference fused with semantic validation - constitutes the principal contribution of this work, and is detailed formally in Section: Materials and Methods.

### Related Work

#### Network-Based Methods and Guilt-by-Association

The earliest computational repurposing methodologies were rooted in the “guilt-by-association” (GBA) heuristic, premised on the assumption that entities occupying proximal or topologically similar positions within a biological network are likely to share functional or phenotypic properties [11]. Early implementations operationalized this principle through network propagation algorithms, most notably the Random Walk with Restart (RWR) formulation:  $p^{(t+1)} = (1 - \alpha) * W * p^{(t)} + \alpha * p^{(0)}$ , where  $W$  denotes the row-normalized adjacency matrix of the interaction network,  $p^{(0)}$  the initial seed distribution over query nodes (e.g., disease-associated genes), and  $\alpha$  the restart probability governing locality of diffusion [12]. Variants of this diffusion formalism underpinned prominent tools such as DIAMOND for disease-module identification and pathway-proximity scoring frameworks that quantified the shortest-path or diffusion-state distance between a candidate drug's target set and a disease's associated gene module [12].

While computationally tractable and biologically interpretable, GBA-diffusion approaches suffer acutely from network sparsity and incompleteness. Curated interactomes such as early protein-protein interaction (PPI) maps capture only a partial, literature-biased slice of the true biological interaction space, and diffusion-based scoring degrades sharply when seed nodes are weakly connected or when true functional paths traverse edge types absent from the observed network [13]. Furthermore, these methods lack any parametric learning capacity -

similarity scores are fixed algebraic functions of network topology, incapable of learning task-adaptive representations from labeled therapeutic outcomes, and are fundamentally unable to incorporate heterogeneous node or edge attributes beyond binary connectivity [14].

### Deep Learning and Graph Neural Networks

The advent of Graph Neural Networks (GNNs) marked a substantive shift from fixed algebraic propagation toward learned, parametric representation functions. The Graph Convolutional Network (GCN) formalized spectral graph convolution as a first-order polynomial approximation of the graph Laplacian, yielding the layer-wise propagation rule:  $H^{(l+1)} = \text{sigma}(D^{-1/2} * A * D^{-1/2} * H^{(l)} * W^{(l)})$ , where  $A = A + I$  is the self-looped adjacency matrix and  $D$  its corresponding degree matrix [15]. Subsequently, Graph Attention Networks (GAT) relaxed the fixed-weight aggregation of GCN by introducing learned attention coefficients over neighborhood edges, allowing the model to adaptively weight neighbor contributions rather than relying on degree-normalized uniform averaging [16]. Both architectures, however, were formulated for homogeneous graphs and operate under transductive constraints.

This homogeneous-transductive coupling proved theoretically inadequate for biomedical applications for two compounding reasons. First, biomedical graphs are irreducibly multi-relational: collapsing drug-target, gene-disease, and drug-drug-interaction edges into a single relation type discards the very inductive bias - relation-specific propagation - that carries biological signal [17]. This motivated the development of multi-relational and heterogeneous GNN variants, such as Relational Graph Convolutional Networks (R-GCN), which parameterize a distinct transformation matrix  $W_r$  per relation type  $r$ :  $h_i^{(l+1)} = \text{sigma}(\sum_r (\sum_j (1/c) * W_r * h_j) + W_0 * h_i)$  [18]. Comparative studies have consistently demonstrated that relation-aware architectures outperform homogeneous baselines by explicit preservation of edge semantics [19].

Second, and independently, the transductive constraint shared by GCN, GAT, and R-GCN precludes generalization to nodes absent at training time - a critical failure mode given the continuous influx of newly characterized compounds. This limitation directly motivated inductive frameworks, most notably GraphSAGE, which reformulates representation learning as a learned aggregation function over sampled local neighborhoods rather than a fixed embedding lookup, thereby enabling forward-pass inference on previously unseen nodes without retraining [20].

### Data and Preprocessing

#### The Stanford BioSNAP Heterogeneous Information Network

The empirical foundation of DeepCure Pro is constructed upon the Stanford BioSNAP collection, a curated repository of biomedical interaction networks aggregated from experimentally validated sources including DrugBank, STRING, and disease-phenotype ontologies [21]. Rather than treating this resource as an isolated bipartite drug-disease matrix, this work integrates its constituent sub-networks into a unified Heterogeneous Information Network (HIN), formally defined as  $G = (V, E)$  with node-type mapping  $\phi: V \rightarrow \{\text{Drug, Protein, Disease}\}$  and edge-type mapping  $\psi: E \rightarrow \{\text{targets,}$

interacts-with, associated-with}. This tripartite construction is essential to preserving the mechanistic chain of biological causality - from chemical structure, through proteomic interaction, to phenotypic disease manifestation - that a flattened, single-relation graph would otherwise obscure.

The raw network is assembled from four discrete relational data sources, each ingested as an independent CSV file prior to unification, as summarized in Table 1.

| File                     | Primary Entities | Key Fields                                              | Approx. Records |
|--------------------------|------------------|---------------------------------------------------------|-----------------|
| biosnap_drugs.csv        | Drug (Compound)  | drug_id, drug_name, SMILES, ATC_code                    | ~5,000          |
| biosnap_interactions.csv | Drug-Protein     | drug_id, protein_id, binding_affinity, interaction_type | ~15,000+        |
| biosnap_ppi.csv          | Protein-Protein  | protein_id_a, protein_id_b, confidence_score            | ~715,000+       |
| biosnap_diseases.csv     | Disease-Gene     | disease_id, disease_name, associated_gene_id, MeSH_code | ~500            |

**Table 1:** Summary of Relational Data Sources Used in Network Construction

The biosnap\_drugs.csv file supplies the canonical SMILES string for each compound, which serves as the raw chemical substrate for downstream molecular featurization. The biosnap\_interactions.csv file encodes empirically validated drug-target binding relationships, forming the targets edge type. The biosnap\_ppi.csv file, by substantial margin the largest of the four sources, encodes the human protein-protein interactome, forming the interacts-with edge type and providing the dense relational substrate through which perturbation signals propagate between drug targets and disease-associated genes. Finally,

biosnap\_diseases.csv supplies gene-phenotype associations curated against Medical Subject Headings (MeSH) taxonomy, forming the associated-with edge type that anchors the graph's disease nodes.

## Graph Construction Logic

### Node Specification and Initialization

Three heterogeneous node populations are instantiated within the unified graph object, each requiring a domain-appropriate feature initialization strategy prior to entry into the neural aggregation pipeline, as summarized in Table 2.

| Node Type | Approx. Count | Raw Feature Source                      | Initial Feature Dimension |
|-----------|---------------|-----------------------------------------|---------------------------|
| Drug      | ~5,000        | SMILES $\rightarrow$ Morgan Fingerprint | 1,024                     |
| Protein   | ~21,000       | Amino-acid sequence embedding           | 512                       |
| Disease   | ~500          | MeSH / Phenotype ontology embedding     | 256                       |

**Table 2:** Node Types and Feature Initialization Strategies

Drug nodes ( $N_{\text{drug}} \approx 5,000$ ) are initialized not through arbitrary learned embedding tables - which would reintroduce the transductive cold-start weakness discussed in Section: The Cold-Start Problem - but through a deterministic, structure-derived molecular fingerprint. Specifically, each drug's canonical SMILES string is algorithmically converted into a 1024-bit Morgan (Extended-Connectivity) Fingerprint, a circular substructure encoding that captures the local chemical topology within a

bond radius  $r$  around each heavy atom via iterative neighborhood hashing [22]. This is formalized in Equation (1):

$$X_{\text{drug}} = \text{MorganFingerprint}(\text{SMILES}, r = 2) \in \{0, 1\}^{1024} \quad (1)$$

Protein nodes ( $N_{\text{protein}} \approx 21,000$ ) are initialized via pretrained sequence-embedding projections derived from amino-acid composition and known domain annotations, yielding a 512-dimensional raw feature vector per node. Disease nodes ( $N_{\text{disease}} \approx 500$ ) are initialized via

ontology-embedding vectors derived from MeSH hierarchical position and associated phenotypic descriptors, yielding a 256-dimensional raw feature vector.

### Hub-Aware Graph Pruning

As established in Section: Hub Bias - The Super-Node Problem, the raw BioSNAP interactome exhibits pronounced scale-free degree distribution, wherein a minority of disease nodes - typically broad, non-specific phenotypic categories - accumulate disproportionately dense connectivity to the gene/protein layer. Left unpruned, these super-nodes dominate the neighborhood-sampling distribution during message passing, biasing the model toward trivial “hub-adjacency” predictions rather than mechanistically specific drug-disease associations. To mitigate this, DeepCure Pro applies a degree-centrality pruning filter at the disease-node level prior to model ingestion. Let  $\deg(v)$  denote the raw connectivity degree of disease node  $v$  within the associated-with edge set; the retained disease vertex set  $V'_{\text{disease}}$  is defined in Equation (2):

$$V'_{\text{disease}} = \{v \in V_{\text{disease}} \mid \deg(v) < \tau_{\text{hub}}, \tau_{\text{hub}} = 500\} \quad (2)$$

Disease nodes exceeding the empirically calibrated threshold  $\tau_{\text{hub}} = 500$  connections are excluded from the training and inference graph, on the interpretive grounds that such nodes represent overly generic phenotypic categories whose connectivity pattern reflects curatorial breadth rather than specific mechanistic association. This pruning step is empirically shown in Section: Experiments and Results to materially reduce false-positive predictions concentrated around high-degree disease categories, directly addressing the hub-bias limitation identified in Section: Introduction.

## MATERIALS AND METHODS

### Heterogeneous GraphSAGE Architecture

#### Feature Projection Layer

Because the three node types enter the pipeline with heterogeneous raw feature dimensionalities (1024 for drugs, 512 for proteins, 256 for diseases, per Table 2), a type-specific linear projection layer is first applied to map all node representations into a shared latent space of unified dimensionality  $d = 128$ . For each node type  $\tau \in \{\text{Drug}, \text{Protein}, \text{Disease}\}$ , an independent projection matrix  $W_{\tau}^{\text{proj}} \in \mathbb{R}^{(d \times d_{\tau})}$  transforms the raw feature vector  $x_v \in \mathbb{R}^{(d_{\tau})}$  of node  $v$  into its initial hidden representation, as given in Equation (3):

$$h_v^{(0)} = \sigma(W_{\phi(v)}^{\text{proj}} x_v + b_{\phi(v)}^{\text{proj}}), h_v^{(0)} \in \mathbb{R}^{128} \quad (3)$$

#### Heterogeneous Message Passing

Following projection, DeepCure Pro performs  $L = 2$  layers of relation-specific SAGE aggregation, implemented

via PyTorch Geometric's HeteroConv container, which internally dispatches a distinct SAGEConv operator per edge type  $r$  in  $R = \{\text{targets}, \text{interacts-with}, \text{associated-with}\}$  and subsequently merges the resulting per-relation messages through summation aggregation at the destination node. Formally, the layer-wise update for node  $v$  of type  $\phi(v)$  at layer  $l+1$  is given by Equation (4):

$$h_v^{(l+1)} = \sigma\left(\sum_{r \in R(\tau)} W_r^{(l)} \cdot \text{MEAN}(\{h_u^{(l)} : u \in N_r(v)\}) + W_{\text{self}}^{(l)} h_v^{(l)}\right) \quad (4)$$

where  $R(\tau)$  denotes the subset of relation types incident to node type  $\tau$ ,  $N_r(v)$  the neighbor set of  $v$  under relation  $r$  (obtained via fixed-size neighborhood sampling rather than full-neighborhood traversal),  $W_r^{(l)}$  a relation-specific learnable weight matrix distinct for each edge type at layer  $l$ , and  $W_{\text{self}}^{(l)}$  a separate self-loop transformation preserving the node's own representation across layers.

### Link Prediction and Hard Negative Mining

#### Similarity-Based Scoring

Following the two-layer heterogeneous encoding pass, each drug node  $u$  and disease node  $v$  possesses a final 128-dimensional latent embedding,  $h_u^{(2)}$  and  $h_v^{(2)}$  respectively. The predicted likelihood of a therapeutic association between  $u$  and  $v$  is computed via L2-normalized cosine similarity over the learned embedding space, as shown in Equation (5):

$$q_{uv} = \text{sim}(u, v) = (h_u^{(2)} \cdot h_v^{(2)}) / (\|h_u^{(2)}\|_2 \|h_v^{(2)}\|_2) \quad (5)$$

#### Weighted Loss with Hard Negative Mining

Because the true therapeutic edge set  $E_{\text{treats}}$  constitutes a minute fraction of all possible drug-disease node pairs, naive random negative sampling produces a training signal dominated by trivially distinguishable negative pairs, yielding poor discriminative resolution near the boundary. DeepCure Pro instead employs a 1:3 hard negative sampling ratio, wherein for every one true positive edge, three negative pairs are drawn preferentially from drug-disease pairs sharing at least one intermediate protein-neighborhood path - i.e., topologically plausible but labeled-negative pairs - forcing the model to learn fine-grained discriminative features rather than exploiting gross topological disconnection as a shortcut.

The model is optimized against a weighted Binary Cross-Entropy (BCE) objective, where the positive class weight  $w_p$  compensates for the residual class imbalance introduced by the 1:3 ratio, as given in Equation (6):

$$L_{\text{BCE}} = -(1/N) \sum_{i=1..N} [w_p \cdot y_i \log(q_i) + (1-y_i) \log(1-q_i)] \quad (6)$$

#### Generative AI RAG Pipeline for Explainability

To resolve the interpretability gap identified in Section: The Interpretability Gap, DeepCure Pro appends a post-

hoc, Retrieval-Augmented Generation (RAG) reasoning layer atop the raw GraphSAGE similarity score. Rather than presenting a stakeholder with an unqualified numerical output  $u_{uv}$ , this layer synthesizes a natural-language, literature-grounded mechanistic rationale by orchestrating Gemini 1.5 Pro against real-time bibliographic retrieval via the Semantic Scholar API. The complete inference-time execution sequence is formalized in Algorithm 1.

#### ALGORITHM 1 RAG-Enhanced Inference Pipeline (DeepCure Pro)

```

1:  $h_u \leftarrow M.encode(u)$  // Eq.(3)-(4) forward pass
2:  $h_v \leftarrow M.encode(v)$ 
3:  $u_{uv} \leftarrow \text{CosineSimilarity}(h_u, h_v)$  // Eq.(5)
4: if  $u_{uv} < \theta$  then
5:
6: return Explained_Prediction(score =  $u_{uv}$ , rationale =
  "Below confidence threshold", citations =  $\emptyset$ )
7: end if
8:
9: path_evidence  $\leftarrow \text{ExtractShortestPaths}(G, u, v, \text{max\_hops} = 3)$ 
10: query_text  $\leftarrow \text{BuildSemanticQuery}(u.name, v.name, \text{path\_evidence})$ 
11:
12: try:
13: lit_results  $\leftarrow \text{SemanticScholarAPI.search}(\text{query\_text}, \text{top\_k} = K)$ 
14: catch APIFailureException:
15: lit_results  $\leftarrow \text{LocalSQLFallback.query}(u.name, v.name)$  // cached corpus
16: end try
17:
18: context_bundle  $\leftarrow \text{Concatenate}(\text{path\_evidence}, \text{lit\_results})$ 
19: prompt  $\leftarrow \text{ConstructGroundedPrompt}(u, v, u_{uv}, \text{context\_bundle})$ 
20:
21: rationale_text  $\leftarrow \text{Gemini1.5Pro.generate}(\text{prompt})$  // RAG synthesis
22: citations  $\leftarrow \text{ExtractCitedSources}(\text{rationale\_text}, \text{lit\_results})$ 
23:
24: return Explained_Prediction(score =  $u_{uv}$ , rationale =
  rationale_text, citations = citations)

```

The pipeline proceeds in three conceptual phases. First (lines 1-7), the trained heterogeneous GraphSAGE model produces the raw structural similarity score per Equations(3)-(5); predictions falling below the operational confidence threshold  $\theta$  are terminated early without incurring generative inference cost. Second (lines 9-17), the algorithm extracts the shortest intermediate mechanistic path(s) connecting drug  $u$  and disease  $v$  through the protein interaction layer, and issues a structured retrieval query against the Semantic Scholar API to surface externally published, peer-reviewed evidence relevant to the candidate association; a local SQL-cached fallback corpus is invoked under API failure or rate-limiting conditions to preserve pipeline availability. Third (lines 18-24), the retrieved literature context is concatenated with the extracted graph path evidence into a single grounding context bundle, which is injected into a structured prompt submitted to Gemini 1.5 Pro for natural-language rationale synthesis, with explicit citation back-attribution to the retrieved source set. This architecture ensures that every DeepCure Pro prediction surfaced to an end user is accompanied not merely by a numerical confidence score, but by a mechanistically traceable, externally verifiable rationale - directly resolving the black-box deficiency articulated in Section: The Interpretability Gap.

### Polypharmacology Synergy Modeling

#### The "Virtual Cocktail" Hypothesis

Beyond single-agent repurposing, a substantial proportion of complex and treatment-resistant pathologies - notably oncological, neurodegenerative, and polygenic metabolic disorders - exhibit therapeutic responses that are more effectively modulated through combinatorial pharmacology than through monotherapy, owing to the multi-pathway, redundant-signaling architecture of the underlying disease network [23]. Empirically screening the combinatorial space of  $n$  approved compounds for synergistic pairs is computationally prohibitive, scaling as  $C(n, 2)$  candidate combinations, and remains experimentally infeasible at the scale of the full BioSNAP drug vocabulary ( $n \approx 5,000$ , yielding over 12 million pairwise candidates).

DeepCure Pro addresses this combinatorial explosion by introducing the "Virtual Cocktail" hypothesis: the proposition that if the learned GraphSAGE latent space is sufficiently well-structured - that is, if geometric proximity within the embedding space is a reliable proxy for shared or complementary mechanistic action - then the therapeutic profile of a hypothetical multi-drug combination can be approximated directly through algebraic composition of the individual drugs' latent embeddings, without requiring an explicit joint-representation retraining pass or combinatorial re-simulation of the full graph.

## Latent Space Arithmetic

Formally, given two candidate drug nodes A and B with converged, L2-normalized final-layer embeddings  $\hat{h}_A^{(2)}$  and  $\hat{h}_B^{(2)}$  (obtained by normalizing the outputs of Equation 4 as in Equation 5), the virtual cocktail centroid vector  $Z_{AB}$ , representing the composite latent signature of the hypothetical combination therapy, is computed as the linear midpoint of the two normalized embeddings, as given in Equation (7):

$$Z_{AB} = (\hat{h}_A^{(2)} + \hat{h}_B^{(2)}) / 2, \hat{h}_i^{(2)} = h_i^{(2)} / \|h_i^{(2)}\|_2 \quad (7)$$

This centroid vector  $Z_{AB} \in \mathbb{R}^{128}$  occupies the same latent manifold as individual drug and disease embeddings, and can therefore be directly substituted into the cosine similarity scoring function of Equation (5) - evaluated against candidate disease embeddings  $h_v^{(2)}$  - to estimate the composite therapeutic alignment of the drug pair against a given indication, as given in Equation (8):

$$u_{AB,v} = (Z_{AB} \cdot h_v^{(2)}) / (\|Z_{AB}\|_2 \|h_v^{(2)}\|_2) \quad (8)$$

The interpretive value of this formulation is twofold.

First, when  $u_{AB,v} > \max(u_{A,v}, u_{B,v})$  - that is, the composite score exceeds both constituent single-drug scores - the pairing is flagged as a candidate for synergistic complementarity, suggesting that the two compounds address non-overlapping regions of the disease-relevant subgraph. Second, when  $Z_{AB}$  drifts toward regions of the embedding space associated with known drug-drug interaction (DDI) adverse-event clusters, the combination is flagged for antagonism or toxicity risk review. Consistent with the design philosophy, any elevated synergy score triggers the same Algorithm 1 pipeline.

## Implementation Challenges and Solutions

### Challenge 1: Super-Node Bias in Message Aggregation

**Problem:** As established in Sections Hub Bias The Super-Node Problem and Hub-Aware Graph Pruning, the raw BioSNAP interactome's scale-free degree distribution causes a small subset of high-connectivity disease and protein nodes to dominate the neighborhood aggregation signal in Equation (4), suppressing gradient contributions from mechanistically specific, low-degree edges and inflating false-positive prediction rates for generic hub-adjacent diseases.

**Solution:** A centrality long-tail pruning threshold is enforced at the graph-construction stage, formalized previously in Equation (2), excluding any disease node with degree centrality  $\deg(v) \geq \tau_{\text{hub}} = 500$  from the training and inference graph. This threshold was calibrated empirically via ablation across candidate values  $\tau_{\text{hub}} \in \{100, 250, 500, 1000\}$ , with  $\tau_{\text{hub}} = 500$  selected as the value minimizing validation-set false-positive rate without materially reducing

recall on the retained disease vocabulary (quantified in Section: Experiments and Results).

### Challenge 2: Multi-Modal Feature Alignment

**Problem:** The three node types populating G originate from structurally incompatible feature domains - binary chemical fingerprints (drugs), continuous sequence embeddings (proteins), and ontology-derived phenotype vectors (diseases) - each of differing raw dimensionality (1024, 512, and 256 respectively, per Table 2). Naively concatenating or directly summing these heterogeneous raw vectors within a shared message-passing layer is mathematically ill-posed, as it conflates dimensionally and distributionally non-commensurable feature spaces.

**Solution:** As formalized in Equation (3), DeepCure Pro instantiates independent, type-specific linear projection matrices -  $W_{\text{drug}^{\text{proj}}} \in \mathbb{R}^{(128 \times 1024)}$ ,  $W_{\text{protein}^{\text{proj}}} \in \mathbb{R}^{(128 \times 512)}$ , and  $W_{\text{disease}^{\text{proj}}} \in \mathbb{R}^{(128 \times 256)}$  - each learned independently during training, mapping every node type into the shared  $d=128$  latent space prior to any cross-type message passing. This decouples the feature alignment problem from the relational aggregation problem: by the time Equation (4) executes, all node representations are already dimensionally and (through joint end-to-end training) distributionally aligned, permitting valid summation-based fusion of relation-specific messages.

### Challenge 3: API Latency and Reliability

**Problem:** The RAG explainability pipeline described in Algorithm 1 depends on two external network-bound services - the Semantic Scholar API for literature retrieval and the Gemini 1.5 Pro endpoint for narrative synthesis. Under production load, both services are subject to intermittent latency spikes, rate-limiting, and transient unavailability, any of which, if unhandled, would cause synchronous blocking of the inference pipeline and degrade end-to-end system responsiveness below acceptable interactive thresholds.

**Solution:** A Circuit Breaker pattern is implemented at the API-orchestration layer, wrapping all external calls with an explicit timeout boundary of 5 seconds. Should the Semantic Scholar API fail to return a valid response within this window (as reflected in Algorithm 1, lines 12-17), the circuit breaker trips and routes the retrieval request to a locally cached SQL/JSON fallback corpus - a periodically refreshed snapshot of previously retrieved literature results indexed by drug-disease query pairs. This ensures that the explainability layer degrades gracefully to a cached-evidence mode under external service disruption, rather than failing the prediction request outright, preserving system availability while transparently flagging to the end user when a rationale has been generated from cached rather than live-retrieved evidence.

### Challenge 4: Heterogeneous Device Management

**Problem:** During development and deployment across heterogeneous compute environments - local CPU-only development machines versus CUDA-enabled GPU inference servers - the framework repeatedly encountered cross-device tensor runtime errors, arising when a subset of model parameters or intermediate activation tensors resided on a GPU device context while downstream operations (e.g., the RAG pipeline's embedding-extraction step, or CPU-bound RDKit fingerprint computation in Equation 1) expected CPU-resident tensors, or vice versa.

**Solution:** A dynamic hardware casting abstraction layer was introduced at the model-serving boundary, wrapping all tensor-producing and tensor-consuming operations with an automatic device-context resolver that inspects the runtime environment (`torch.cuda.is_available()`) at initialization and transparently casts all model weights, input tensors, and intermediate embeddings to a single consistently resolved device (`torch.device`) prior to any cross-module hand-off. This abstraction guarantees that identical inference code executes correctly and deterministically regardless of the underlying host hardware configuration, materially improving deployment portability between development and production environments.

### System Architecture and Implementation

#### Backend Infrastructure

| Endpoint              | Method | Function                                                                                                                                             |
|-----------------------|--------|------------------------------------------------------------------------------------------------------------------------------------------------------|
| /predict/top_diseases | POST   | Given a drug identifier, returns the top-k ranked candidate disease indications using Equation (5), each accompanied by a RAG-generated explanation. |
| /analyze_drug         | GET    | Returns the complete node profile, latent embedding summary, and structural metadata for a selected drug.                                            |
| /analyze_combination  | POST   | Given two drug identifiers, computes the Virtual Cocktail centroid Equation (7) and predicts disease alignment using Equation (8).                   |

**Table 3:** RESTful API Endpoints and Their Functional Descriptions

The `/predict/top_diseases` endpoint constitutes the primary single-agent repurposing interface, internally invoking the full Algorithm 1 pipeline per candidate disease. The `/analyze_drug` endpoint supports exploratory inspection of a compound's learned representation independent of a specific disease query. The `/analyze_combination` endpoint operationalizes the polypharmacology synergy formulation of Section: Polypharmacology Synergy Modeling, exposing the latent-arithmetic centroid computation as a directly queryable service.

The DeepCure Pro inference service is implemented atop FastAPI, selected for its native support of asynchronous request handling - a design requirement directly motivated by the latency-sensitive, I/O-bound external API dependencies discussed in Section: Challenge 3: API Latency and Reliability. Rather than blocking a worker process for the full duration of a Semantic Scholar or Gemini 1.5 Pro round-trip, the backend leverages Python's `async/await` concurrency primitives to permit the server to concurrently service multiple in-flight prediction requests while individual RAG sub-calls await external I/O completion, substantially improving aggregate throughput under concurrent user load relative to a synchronous WSGI implementation.

To eliminate repeated disk I/O and model deserialization latency on a per-request basis, the trained heterogeneous GraphSAGE model - comprising the projection matrices of Equation (3), the relation-specific aggregation weights of Equation (4), and all associated PyTorch Geometric graph-state buffers - is loaded once into server memory at application startup and persisted for the lifetime of the running process, with subsequent inference requests operating directly against the resident in-memory model object rather than triggering reload cycles.

The backend exposes its predictive and explainability functionality through three principal RESTful endpoints, summarized in Table 3.

#### Frontend Interface

The user-facing layer of DeepCure Pro is implemented as an interactive Streamlit application, chosen for its capacity to rapidly compose data-driven dashboard components without a dedicated JavaScript frontend build pipeline, while retaining sufficient interactivity for exploratory pharmacological analysis. The interface is organized into three coordinated panels: a drug/disease query panel, permitting free-text or autocomplete-assisted selection of compounds and indications from the pruned BioSNAP vocabulary (Section Hub-Aware Graph Pruning); a ranked prediction dashboard, rendering the top-k disease associations returned by `/predict/top_diseases` as an interactive, sortable table annotated

with confidence scores and expandable RAG-generated rationale text with inline Semantic Scholar citation links; and a molecular visualization panel, which leverages RDKit's 2D depiction engine to render the queried compound's structural diagram directly from its SMILES representation, providing the end user with immediate visual chemical context alongside the network-derived predictive output. This visualization panel additionally renders side-by-side structural comparisons for the two constituent compounds when the user queries the `/analyze_combination` endpoint, allowing direct visual inspection of the chemical basis underlying a proposed virtual cocktail hypothesis.

## EXPERIMENTS AND RESULTS

### Experimental Setup

All model development, training, and inference benchmarking were conducted within a controlled software environment to ensure reproducibility of the reported results. The heterogeneous GraphSAGE architecture described in Section: Materials and Methods was implemented using PyTorch Geometric (PyG) version 2.3.0, selected specifically for its native HeteroConv and SAGEConv operator support underpinning Equations (3)-(4). Molecular fingerprint extraction (Equation 1) was performed using RDKit version 2022.09, operating directly on the canonical SMILES strings sourced from biosnap\_drugs.csv. All training runs were executed on an NVIDIA GPU execution environment, leveraging CUDA-accelerated sparse tensor operations to accommodate the scale of the pruned BioSNAP heterogeneous graph (Section Hub-Aware Graph Pruning), with the dynamic

hardware-casting abstraction layer (Section Challenge 4: Heterogeneous Device Management) ensuring consistent tensor placement throughout the training loop.

Model parameters were optimized using the Adam optimizer, applied against the weighted Binary Cross-Entropy objective of Equation (6), with the 1:3 hard negative sampling ratio enforced at each mini-batch construction step. The labeled therapeutic edge set was partitioned using an 80-20 train-test split, stratified to preserve the relative proportion of disease categories across both partitions and ensuring that no evaluation-set drug-disease edge was visible to the model during the neighborhood-sampling forward passes of training. Model selection across hyperparameter configurations (learning rate, hidden dimensionality, neighborhood sample size) was performed via a held-out validation subset carved from the training partition, with the final reported test-set metrics computed only once, on the fully held-out 20% partition, to avoid validation-set leakage.

### Overall Performance Comparison

DeepCure Pro was benchmarked against two representative baseline classes: a classical Matrix Factorization approach, representative of pre-neural network-based link prediction (Section Network-Based Methods and Guilt-by-Association), and a Homogeneous Graph Convolutional Network (GCN), representative of relation-agnostic deep learning approaches (Section Deep Learning and Graph Neural Networks) that discard the heterogeneous edge-type structure preserved in Equation (4). Table 4 presents the comparative results.

| Method               | AUC-ROC | Accuracy |
|----------------------|---------|----------|
| Matrix Factorization | 0.852   | 78.50%   |
| Homogeneous GCN      | 0.891   | 82.40%   |
| DeepCure Pro (Ours)  | 0.934   | 88.70%   |

**Table 4:** Comparative Results

The results in Table 4 exhibit a clear, monotonic performance gradient consistent with the theoretical motivation established in Sections Introduction and Related Work. Matrix Factorization, lacking any capacity to incorporate node-level chemical or sequence features and operating purely on observed adjacency structure, yields the weakest AUC-ROC (0.852), corroborating the sparsity limitation discussed in Section Network-Based Methods and Guilt-by-Association. The Homogeneous GCN improves substantially upon this baseline (AUC-ROC 0.891) by introducing learned, feature-aware representation learning, yet remains constrained by its collapse of the targets, interacts-with, and associated-with edge types into an undifferentiated adjacency matrix - precisely the heterogeneity limitation formalized in Section: Data Heterogeneity. DeepCure Pro's relation-aware heterogeneous aggregation (Equation 4), combined

with hub-aware pruning (Equation 2) and hard negative mining (Section: Weighted Loss with Hard Negative Mining), yields a further 4.3-percentage-point AUC-ROC improvement and a 6.3-percentage-point accuracy improvement over the homogeneous GCN baseline, empirically validating the central architectural hypothesis of this work: that explicit preservation of relational edge semantics, combined with degree-bias correction, yields materially superior discriminative capacity on real-world biomedical link prediction.

### System Latency Profile

Beyond predictive accuracy, the operational viability of DeepCure Pro as a deployed clinical-research tool depends critically on end-to-end response latency, particularly given the external API dependencies introduced by the RAG explainability pipeline (Algorithm 1). Table 5 decomposes the full inference pipeline into its constituent

latency contributions, measured under standard operating conditions.

| Operation                                  | Latency |
|--------------------------------------------|---------|
| SMILES → Fingerprint Generation (Eq. 1)    | 8 ms    |
| GraphSAGE Forward-Pass Inference (Eq. 3-5) | 20 ms   |
| Local Database Search                      | 15 ms   |
| External Semantic Scholar API Search       | 2.5 s   |
| Gemini 1.5 Pro Rationale Synthesis         | 3.0 s   |
| Total End-to-End Latency                   | ≈ 5.5 s |

**Table 5:** Inference Pipeline Latency Breakdown Under Standard Operating Conditions

The latency decomposition in Table 5 reveals a pronounced asymmetry between the structural inference path and the semantic explainability path. The core predictive pipeline - fingerprint generation, GraphSAGE forward inference, and local database lookup - collectively completes in under 45 milliseconds, confirming that the heterogeneous GraphSAGE backbone itself imposes negligible computational overhead and is well within the bounds of real-time interactive use. The overwhelming majority of total system latency (over 99% of the roughly 5.5-second total) is attributable to the two external generative and retrieval calls comprising the RAG explainability layer. This asymmetry directly substantiates the architectural necessity of the Circuit Breaker pattern

introduced in Section Challenge 3: API Latency and Reliability: because the explainability layer's latency profile is dominated by network-bound, externally-hosted services rather than the locally-controlled model, graceful degradation to the cached fallback corpus under the 5-second timeout boundary is essential to bounding worst-case user-facing response time.

### Drug Combination Synergy Predictions

To empirically illustrate the polypharmacology synergy formulation of Section: Polypharmacology Synergy Modeling, Table 6 presents representative outputs of the /analyze\_combination endpoint, applying the virtual cocktail centroid computation of Equation (7) and the composite alignment scoring of Equation (8) to two clinically motivated drug pairings.

| Drug A       | Drug B    | Target Disease                                   | Score | Mechanism                                                                                                                                                            |
|--------------|-----------|--------------------------------------------------|-------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Aspirin      | Metformin | Type 2 Diabetes-associated Vascular Inflammation | 0.87  | Complementary anti-inflammatory (COX-mediated) and insulin-sensitizing (AMPK-mediated) pathway convergence on vascular endothelial gene neighbourhoods.              |
| Atorvastatin | Metformin | Metabolic Syndrome                               | 0.91  | Combined lipid-regulatory (HMG-CoA reductase pathway) and glycemic-regulatory (AMPK pathway) target neighbourhoods with shared downstream inflammatory gene overlap. |

**Table 6:** Representative Drug Combination Analysis Results Using the Polypharmacology Synergy Framework

Both illustrative pairings yield composite alignment scores ( $u_{AB,v}$ , Equation 8) exceeding the individual single-agent similarity scores of either constituent compound in isolation, satisfying the synergy criterion articulated in Section: Latent Space Arithmetic ( $u_{AB,v} > \max(u_{A,v}, u_{B,v})$ ).

The RAG explainability layer's mechanistic annotations for both pairings correctly surface non-overlapping pathway involvement - anti-inflammatory versus insulin-sensitizing action in the Aspirin-Metformin case, and lipid-regulatory versus glycemic-regulatory action in the Atorvastatin-Metformin case - consistent with established pharmacological literature on complementary multi-target intervention in metabolic disease, and demonstrating

that the latent-arithmetic centroid formulation of Equation (7) produces combination hypotheses that are not merely numerically elevated but mechanistically coherent upon RAG-based literature grounding.

## CONCLUSION AND FUTURE WORK

This paper presented DeepCure Pro, an end-to-end computational framework for AI-driven drug repurposing that directly confronts the four principal deficiencies constraining prior graph-based approaches: data heterogeneity, hub-induced predictive bias, the interpretability gap of black-box neural inference, and the cold-start limitation of transductive embedding methods. By coupling an inductive, relation-aware heterogeneous GraphSAGE backbone - initialized over chemically-grounded Morgan fingerprint representations and refined through hub-aware graph pruning and hard-negative-mined link prediction - with a Retrieval-Augmented Generation explainability layer orchestrating Gemini 1.5 Pro against live Semantic Scholar literature grounding, DeepCure Pro achieves an AUC-ROC of 0.934 and classification accuracy of 88.7% on the Stanford BioSNAP heterogeneous biomedical network, materially exceeding both classical matrix factorization and homogeneous GCN baselines. Beyond single-agent prediction, the introduction of latent space arithmetic for polypharmacology synergy estimation (Section Polypharmacology Synergy Modeling) demonstrates that the learned embedding manifold supports compositional reasoning over combination therapies without requiring explicit joint-representation retraining, extending the framework's translational utility beyond monotherapy hypothesis generation.

Not with standing these results, three concrete directions are identified for future extension of this work:

### Transcriptomic Integration

The current node feature schema (Table 2) is limited to static structural and sequence-derived representations. Future iterations of DeepCure Pro will incorporate differential gene expression profiles (e.g., LINCS L1000 perturbation signatures) as dynamic, condition-specific node features, enabling the model to distinguish disease-state-dependent transcriptomic perturbation patterns rather than relying solely on static interactome topology - a refinement expected to particularly improve prediction fidelity for diseases with heterogeneous molecular subtypes.

### Attention-Based Architectural Upgrades

While the mean-aggregation formulation of Equation (4) provides strong inductive generalization, it assigns uniform importance to all sampled neighbors within a given relation type. Future work will investigate replacing the relation-specific SAGEConv operators with heterogeneous Graph Attention (GAT) layers, enabling the model to learn adaptive, edge-specific attention coefficients that may

further mitigate residual hub-bias effects beyond the fixed-threshold pruning strategy of Equation (2), while simultaneously providing an additional, attention-weight-derived source of structural interpretability to complement the RAG semantic layer.

### Clinical Trial Outcome Prediction Module

The present framework terminates its predictive scope at the hypothesis-generation stage. A natural extension involves augmenting DeepCure Pro with a downstream clinical trial outcome prediction module, trained on historical ClinicalTrials.gov phase-transition and adverse-event data, to estimate the probability of clinical success for a given computationally repurposed candidate prior to wet-lab or trial-stage investment - thereby extending the framework's contribution beyond hypothesis generation into probabilistic prioritization across the full translational pipeline described in Section: The Economics of Pharmaceutical Failure: Eroom's Law.

## REFERENCES

1. Scannell JW, Blanckley A, Boldon H, Warrington B. Diagnosing the decline in pharmaceutical R&D efficiency. *Nat. Rev. Drug Discov.* 2012;11(3):191-200.
2. DiMasi JA, Grabowski HG, Hansen RW. Innovation in the pharmaceutical industry: new estimates of R&D costs. *J. Health Econ.* 2016;47:20-33.
3. Ashburn TT, Thor KB. Drug repositioning: identifying and developing new uses for existing drugs. *Nat. Rev. Drug Discov.* 2004;3(8):673-683.
4. Pushpakom S, Iorio F, Eyers PA, Escott KJ, Hopper S, Wells A, et al. Drug repurposing: progress, challenges and recommendations. *Nat. Rev. Drug Discov.* 2019;18(1):41-58.
5. Barabási AL, Gulbahce N, Loscalzo J. Network medicine: a network-based approach to human disease. *Nat. Rev. Genet.* 2011;12(1):56-68.
6. Fey M, Lenssen JE. Fast graph representation learning with PyTorch Geometric. *arXiv preprint arXiv:1903.02428.* 2019.
7. Hamilton W, Ying Z, Leskovec J. Inductive representation learning on large graphs. *Adv. Neural Inf. Process. Syst.* 2017;30.
8. Hamilton WL, Ying R, Leskovec J. Representation learning on graphs: Methods and applications. *arXiv preprint arXiv:1709.05584.* 2017.
9. Dong Y, Chawla NV, Swami A. metapath2vec: Scalable representation learning for heterogeneous networks. In *Proceedings of the 23rd ACM SIGKDD international conference on knowledge discovery and data mining.* 2017;135-144.
10. Wang X, Ji H, Shi C, Wang B, Ye Y, Cui P, Yu PS. Heterogeneous graph attention network. In *The world wide web conference.* 2019;2022-2032.
11. Perozzi B, Al-Rfou R, Skiena S. Deepwalk: Online learning of social representations. In *Proceedings of the 20th ACM SIGKDD international conference on Knowledge discovery and data mining* 2014;701-710.
12. Lundberg SM, Lee SI. A unified approach to interpreting model predictions. *Adv. Neural Inf. Process. Syst.* 2017;30.

13. Lewis P, Perez E, Piktus A, Petroni F, Karpukhin V, Goyal N, et al. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Adv. Neural Inf. Process. Syst.* 2020;33:9459-9474.
14. Szklarczyk D, Gable AL, Nastou KC, Lyon D, Kirsch R, Pyysalo S, et al. The STRING database in 2021: customizable protein-protein networks, and functional characterization of user-uploaded gene/measurement sets. *Nucleic Acids Res.* 2021;49(D1):D605-612.
15. Mikolov T, Sutskever I, Chen K, Corrado GS, Dean J. Distributed representations of words and phrases and their compositionality. *Adv. Neural Inf. Process. Syst.* 2013;26.
16. Chou TC. Theoretical basis, experimental design, and computerized simulation of synergism and antagonism in drug combination studies. *Pharmacol. Rev.* 2006;58(3):621-681.
17. Subramanian A, Narayan R, Corsello SM, Peck DD, Natoli TE, Lu X, et al. A next generation connectivity map: L1000 platform and the first 1,000,000 profiles. *Cell.* 2017;171(6):1437-1452.
18. Lamb J, Crawford ED, Peck D, Modell JW, Blat IC, Wrobel MJ, et al. The Connectivity Map: using gene-expression signatures to connect small molecules, genes, and disease. *science.* 2006;313(5795):1929-1935.
19. Veličković P, Cucurull G, Casanova A, Romero A, Lio P, Bengio Y. Graph attention networks. In *International conference on learning representations* 2018;6(2).
20. Wang X, Bo D, Shi C, Fan S, Ye Y, Yu PS. A survey on heterogeneous graph embedding: methods, techniques, applications and sources. *IEEE transactions on big data.* 2022;9(2):415-436.
21. Ashish V. Attention is all you need. *Adv. Neural Inf. Process. Syst.* 2017;30:1.
22. Zarin DA, Tse T, Williams RJ, Califf RM, Ide NC. The ClinicalTrials.gov results database-update and key issues. *N. Engl. J. Med.* 2011;364(9):852-860.
23. Gayvert KM, Madhukar NS, Elemento O. A data-driven approach to predicting successes and failures of clinical trials. *Cell Chem. Biol.* 2016;23(10):1294-1301.
24. Kuhn M, Szklarczyk D, Pletscher-Frankild S, Blicher TH, Von Mering C, Jensen LJ, Bork P. STITCH 4: integration of protein-chemical interactions with user data. *Nucleic Acids Res.* 2014;42(D1):D401-D407.
25. Wishart DS, Feunang YD, Guo AC, Lo EJ, Marcu A, Grant JR, et al. DrugBank 5.0: a major update to the DrugBank database for 2018. *Nucleic Acids Res.* 2018;46(D1):D1074-D1082.
26. J. Piñero et al. The DisGeNET knowledge platform for disease genomics: 2019 update. *Nucleic Acids Res.* 2020;48(D1):D845-D855.
27. Kipf TN, Welling M. Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:1609.02907.* 2016.
28. Zitnik M, Agrawal M, Leskovec J. Modeling polypharmacy side effects with graph convolutional networks. *Bioinformatics.* 2018;34(13):i457- i466.
29. D. Rogers and M. Hahn. Extended-connectivity fingerprints. *J. Chem. Inf. Model.* 2010;50(5):742-754.
30. G. Landrum. RDKit: Open-source cheminformatics software. *RDKit Documentation.* 2022.
31. Santra AK, Christy CJ. Genetic algorithm and confusion matrix for document clustering. *Int. J. Comput. Sci. Issues.* 2012;9(1):322.
32. Schlichtkrull M, Kipf TN, Bloem P, Van Den Berg R, Titov I, Welling M. Modeling relational data with graph convolutional networks. In *European semantic web conference.* 2018;593-607. Cham: Springer International Publishing.
33. Luo Y, Zhao X, Zhou J, Yang J, Zhang Y, Kuang W, et al. A network integration approach for drug-target interaction prediction and computational drug repositioning from heterogeneous information. *Nature communications.* 2017;8(1):573.
34. Wan F, Hong L, Xiao A, Jiang T, Zeng J. NeoDTI: neural integration of neighbor information from a heterogeneous network for discovering new drug-target interactions. *Bioinformatics.* 2019;35(1):104-111.
35. Gawehn E, Hiss JA, Schneider G. Deep learning in drug discovery. *Molecular informatics.* 2016;35(1):3-14.