

M3NetFlow: A novel multi-scale multi-hop graph AI model for integrative multi-omic data analysis

Heming Zhang¹, S. Peter Goedegebuure^{2,3}, Li Ding^{3,4}, David DeNardo^{3,4}, Ryan C. Fields^{2,3}, Yixin Chen⁶, Philip Payne¹, Fuhai Li^{1,5#}

¹Institute for Informatics, Data Science and Biostatistics (I2DB), ²Department of Surgery,

³Siteman Cancer Center, ⁴Department of medicine, ⁵Department of Pediatrics, Washington

University School of Medicine, ⁶Department of Computer Science and Engineering, Washington

University in St. Louis, St. Louis, MO, USA. #Correspondence: Fuhai.Li@wustl.edu

Summary: Multi-omic data-driven studies, characterizing complex disease signaling system from multiple levels, are at the forefront of precision medicine and healthcare. The integration and interpretation of multi-omic data are essential for identifying molecular targets and deciphering core signaling pathways of complex diseases. However, it remains an open problem due the large number of biomarkers and complex interactions among them. In this study, we propose a novel Multi-scale Multi-hop Multi-omic graph model, *M3NetFlow*, to facilitate generic multi-omic data analysis to rank targets and infer core signaling flows/pathways. To evaluate *M3NetFlow*, we applied it in two independent multi-omic case studies: 1) uncovering mechanisms of synergistic drug combination response (defined as anchor-target guided learning), and 2) identifying biomarkers and pathways of Alzheimer's disease (AD). The evaluation and comparison results showed *M3NetFlow* achieves the best prediction accuracy (accurate), and identifies a set of essential targets and core signaling pathways (interpretable). The model can be directly applied to other multi-omic data-driven studies. The code is publicly accessible at: <https://github.com/FuhaiLiAiLab/M3NetFlow>

1. Introduction

Multi-omic data-driven studies are at the forefront of precision medicine and healthcare. Recently, multi-omic datasets, like genetic, epigenetic, transcriptomic, and proteomic, have been being generated to characterize dysfunctional biological processes and signaling pathways from multiple levels/views, and to elucidate the panoramic view of the disease pathogenesis¹⁻⁵. For example, The Cancer Genome Atlas (TCGA) program have generated multi-omic datasets of over 20,000 samples spanning 33 cancer types, to understand the key molecular targets and signaling pathways of cancer⁶. Moreover, the multi-omic data of >10,000 cancer cell lines were profiled in the Cancer Cell Line Encyclopedia (CCLE) project, which are valuable to investigate the mechanism of cancer response to given drugs and drug combinations⁷. In addition, the multi-omic data of Alzheimer's disease (AD) are generated and publicly available in the ROSMAP⁸ project to uncover the pathogenesis of AD. Also, the exceptional longevity (EL), like the Long-Life Family Study (LLFS) project, have been generating multi-omic data⁹⁻¹¹ to identify protective biomarkers and pathways for long and healthy life. The multi-omic data are valuable and essential for understanding the key molecular targets and mechanisms of diseases, identifying novel therapeutic targets, predicting effect drugs and drug cocktails to guide the development of precision medicine.

However, it remains an open problem and a challenging task for integrative and interpretable omic data analysis, though a set of computational models has been proposed. A comprehensive review of existing multi-omic data integration analysis models was reported¹². Specifically, these models were clustered into a few categories, like similarity, correlation, Bayesian, multivariate, fusion and network-based models. The PAtchway Representation and Analysis by Direct Inference on Graphical Models (PARADIGM¹³) is one of the most widely used methods among these traditional computational methods. Aside from that, the NMTF-based method called iCell¹⁴ was also

proposed to integrate the multi-omic data and only the direct connection was considered. The timeOmics¹⁵ was recently developed to identify signaling patterns over time of biomarkers in multi-omic data, allowing for the analysis of time-series data in biological systems.

Problem formulation. The problem to be tackled in this study using graph AI models is to identify or rank disease or drug response associated molecular biomarkers from a large number of multi-omic features, and infer the potential signaling network among the selected biomarkers. The two specific biological problems to be tackled using the graph AI models are as follows. The first task is an anchor-biomarker (like known molecular targets, or drug targets) guided learning, e.g., to identify targets and pathways regulating drug cocktail response by integrating multi-omic data of cells and drug targets. The input data of the first case study is the multi-omic profiling of individual cell lines, kegg signaling pathways (graph), drug targets, and synergistic scores of a set of drug combinations. The output of the model is a set of top-ranked biomarkers and signaling pathways linking to the given anchor-biomarkers (information flowing to the anchor-biomarkers), which are informative to predict and explain the mechanism of drug combination response. In addition to drug-targets, the anchor-biomarkers can be generic given targets of interested to study the up- or down-stream signaling pathways of the anchor-biomarkers. The second task is to identify disease associated biomarkers and signaling pathways of Alzheimer's disease (AD). It is a generic biomarker ranking and pathway inference problem without the guide/constraint of anchor-biomarkers. The input is multi-omic data, like the AD and control/normal samples and kegg signaling pathways (graph). The output of the model is a set of top-ranked biomarkers and signaling pathways, which are informative to classify AD from normal samples explaining the pathogenesis of AD. The two biological problems/case studies represent two of the most needed multi-omic data analysis tasks.

Related work. Graph Neural Networks (GNNs) have gained prominence due to their capability

to model relationships within graph-structured data^{16–19}. And numerous studies have applied the Graph Neural Network (GNN) with the integration of the multi-omic data. MOGONET²⁰ initially creates similarity graphs among samples by leveraging each omic data, then employs a Graph Convolutional Network (GCN¹⁶) to learn a label distribution from each omic data independently. Subsequently, a Cross-omic discovery tensor is implemented to refine the prediction by learning the dependency among multi-omic data. MoGCN²¹ adopts a similar approach by constructing a patient similarity network using multi-omic data and then using GCN to predict the cancer subtype of patients. GCN-SC²² utilizes a GCN to combine single-cell multi-omic data derived from varying sequencing methodologies. MOGCL²³ takes this further by exploiting the potency of graph contrastive learning to pretrain the GCN on the multi-omic dataset, thereby achieving impressive results in downstream tasks with fine-tuning. Nevertheless, none of the aforementioned techniques contemplate incorporating structured signaling data like KEGG into the model. Moreover, general GNN models are limited by their expression power, i.e., the low-pass filtering or over-smoothing issues, which hampers their ability to incorporate many layers. The over-smoothing problem was firstly mentioned by extending the propagation layers in GCN²⁴. Moreover, theoretical papers using Dirichlet energy showed diminished discriminative power by increasing the propagation layers²⁵. And multiple attempts were made to compare the expressive power of the GCNs^{19,26}, and it is shown that WL subtree kernel²⁷ is insufficient for capturing the graph structure. Hence, to improve the expression powerful of GNN, the K -hop information of local substructure was considered in various recent research^{28–33}. However, none of these studies was specifically designed to well integrate the biological regulatory network and provide the interpretation with important edges and nodes.

In this study, we present a novel graph AI model, named ***M3NetFlow (Multi-scale, Multi-Hop, Multi-omic NETWORK Flow inference)*** to address the challenges mentioned above. Specifically, using an attention mechanism, our approach tackles these challenges by first incorporating local

K-hop information within each subgraph by integrating the gene regulatory network such as KEGG³⁴ and BioGrid^{35,36}. Subsequently, we establish global-level bi-directional propagation to facilitate message passing between genes/proteins. In the interpretation phase, we leverage the attention mechanism to aggregate the weights of all connected paths for genes, and a reweighting process inspired by (Term Frequency – Inverse Document Frequency) TF-IDF³⁷ is employed to redistribute the importance scores of genes across different cell lines. Afterward, a visualization tool, **NetFlowVis**, was developed for the better analysis of targets and signaling pathways of drugs and drug combinations. To assess and demonstrate the effectiveness of our proposed model, **M3NetFlow**, applied it in two independent multi-omic case studies: 1) uncovering mechanisms of synergistic drug combination response (defined as anchor-target guided learning), and 2) identifying biomarkers and pathways of Alzheimer’s disease (AD). The evaluation and comparison results showed *M3NetFlow* achieves the best Pearson correlation or prediction accuracy (accurate), and identifies a set of essential targets and core signaling pathways (interpretable, which can be directly applied to other multi-omic data-driven studies).

2. Methodology

Table 1 provided the links from which four types of datasets, drug combination effects, multi-omic data, gene-gene interactions, and drug-gene interactions, were collected to predict drug combinations for cancer cell lines. What’s more, to investigate Alzheimer’s disease, we also obtained corresponding multi-omic datasets and clinical dataset from publicly available sources, particularly the ROSMAP datasets (see **Table 2**). Comprehensive details regarding the data and the preprocessing results are thoroughly documented in Appendix Section A. This section provides an in-depth explanation of the methodologies employed, as well as the specific parameters and outcomes obtained during the preprocessing stage.

Table 1. Multi-omic data and drug combination screening data of cancer cell lines.

Database Type	Database Name	Website Link
Drug Combination Effect Data	NCI ALMANAC ⁴¹ And O'Neil ⁴² drug combination	https://wiki.nci.nih.gov/display/NCIDTPdata/NCI-ALMANAC and https://drugcomb.fimm.fi/
Multi-omic Data	Cell Model Passports ⁴³ RNA-Seq	https://cog.sanger.ac.uk/cmp/download/rnaseq_20191101.zip
	Cell Model Passports ⁴³ CNV	https://cog.sanger.ac.uk/cmp/download/cnv_20191101.zip
	CCLE ⁴⁴ Methylation	https://data.broadinstitute.org/ccle/CCLE_RRBS_TSS1kb_20181022.txt.gz
	CCLE ⁴⁴ Gene Amplification	https://data.broadinstitute.org/ccle_legacy_data/binary_calls_for_copy_number_and_mutation_data/CCLE_MUT_CNA_AMP_DEL_binary_Revealer.gct
	CCLE ⁴⁴ Gene Deletion	https://data.broadinstitute.org/ccle_legacy_data/binary_calls_for_copy_number_and_mutation_data/CCLE_MUT_CNA_AMP_DEL_binary_Revealer.gct
Gene-Gene Interaction Data	KEGG ⁴⁵ Gene Interactions Network	https://www.genome.jp/kegg/
Drug-Target Interaction Data	DrugBank ⁴⁶	https://go.drugbank.com/

Table 2. Multi-omic data and clinical data from ROSMAP database resources

Database Type	Database Name	Website Link
Clinical Data	ROSMAP_clinical	https://www.synapse.org/#!Synapse:syn3191087
Multi-omic Data	ROSMAP_arrayMethylation_imputed	https://www.synapse.org/#!Synapse:syn3168763
	ROSMAP.CNV.Matrix(Mutation)	https://www.synapse.org/#!Synapse:syn26263118
	ROSMAP_RNAseq_FPKM_gene	https://www.synapse.org/#!Synapse:syn3505720
	C2.median_polish_corrected_log2(Proteomic)	https://www.synapse.org/#!Synapse:syn21266454
Mapping Data	GEO GPL16304 Platform	https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GPL16304

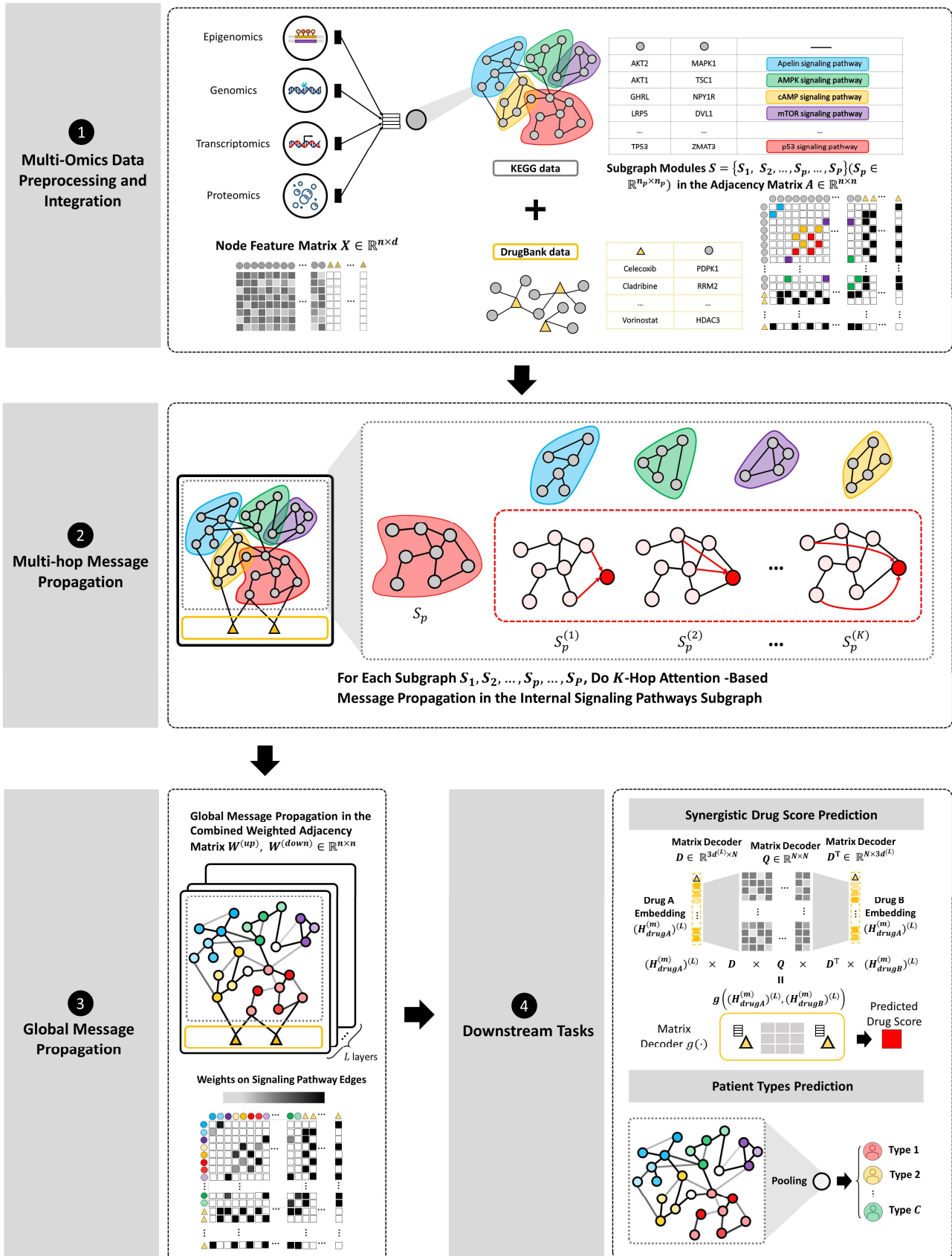


Figure 1. Model architecture of *M3NetFlow*. ① Integrate the one-hot encoded and multi-omic features into a vector for each node. Afterwards, Merge drug-gene and gene-gene interactions into an adjacency matrix. ② Multi-hop attention-based propagation was performed in the subgraphs. ③ Use combined weighted adjacency matrix for global signaling propagation. ④ Downstream tasks. (1) Decode the pair of drug nodes to predict the drug synergy scores. (2) Use pooling strategy to predict the patient outcomes.

2.1 Model Architecture of M3NetFlow

Figure 1 shows the schematic architecture of the proposed *M3NetFlow* model. The model input parameters are: $\mathcal{X} = \{(X^{(1)}, T^{(1)}), (X^{(2)}, T^{(2)}), \dots, (X^{(m)}, T^{(m)}), \dots, (X^{(M)}, T^{(M)})\}$ ($X^{(m)} \in \mathbb{R}^{n \times d}, T \in \mathbb{R}^{n \times 2}$), $A \in \mathbb{R}^{n \times n}$, $S = \{S_1, S_2, \dots, S_p, \dots, S_p\}$ ($S_p \in \mathbb{R}^{n_p \times n_p}$), $D_{in} \in \mathbb{R}^{n \times n}$, $D_{out} \in \mathbb{R}^{n \times n}$, where M represents number of data points in the drug screening dataset. To predict the drug combination effect scores Y ($Y \in \mathbb{R}^{M \times 1}$), we would like to build up the machine learning model $f(\cdot)$ with $f(\mathcal{X}, A, S, D_{in}, D_{out}) = Y$, where $\mathcal{X}$ denotes all of the data points in the dataset and $(X^{(m)}, T^{(m)})$ is m -th data points in the dataset, where $X^{(m)}$ denotes the node features matrix with n nodes of d features and $T^{(m)}$ denotes the one-hot encoding of the number of drugs targeted on those n nodes. The matrix A is the adjacency matrix that demonstrates the node-node interactions, and the element in adjacency matrix A such as a_{ij} indicates an edge from i to j , and S is a set of subgraphs that partition the whole graph adjacent matrix A into multiple subgraphs with $S_p \in \mathbb{R}^{n_p \times n_p}$ of nodes interactions between its internal n_p nodes and each subgraph has its own corresponding subgraph node feature matrix $X_p \in \mathbb{R}^{n_p \times d}$. D_{in} is an in-degree diagonal matrix for nodes in directed graph, and D_{out} is an out-degree diagonal matrix for nodes in directed graph.

Network Modular and Multi-hop Message Propagation In the graph message passing stages of our architecture (see **Figure 1** step3 and step4), the multi-scale design, i.e., the local network module/subgraph module message passing stage for each signaling pathway, and the global message passing stage were designed. The multi-scale design will ensure the message fully interacts in the internal subgraph. Moreover, we opted for a K -hop attention-based graph neural

network because it further allows for the consideration of longer distance information, which can be interpreted as signaling flow in a single message propagation layer. With those initial embedding features $X^{(m)} \in \mathbb{R}^{n \times d}$, ($m = 1, 2, \dots, M$), the K -hop attention-based graph neural network was built to incorporate the long-distance information for each subgraph with

$$\left(\alpha_{ij}^{(m)}\right)_p^{(k)} = \text{ATT}_{hop} \left[X_p^{(m)}\right] \quad (1)$$

$$\left(H_p^{(m)}\right)_i = \text{MSG}_{hop} \left[X_p^{(m)}, \left(\alpha_{ij}^{(m)}\right)_p^{(k)}\right] \quad (2)$$

, where $\left(\alpha_{ij}^{(m)}\right)_p^{(k)}$ is the attention score between node i and node j in the k -th hop of the subgraph S_p for the data point m and the updated node feature $\left(H_p^{(m)}\right)_i$ for node i will be generated via K -hop message propagation (see **Appendix B.1** for details).

Global Bi-directional Message Propagation Following the message propagation in the multiple internal subgraphs, the global weighted bi-directional message propagation will be performed, where nodes-flow contains both ‘upstream-to-downstream’ (from up-stream signaling to drug targets) and ‘downstream-to-upstream’ (from drug targets to down-stream signaling) (see **Figure 1** step 3). Before the global level message propagation, the node feature for data point m in each subgraph $H_p^{(m)}$ ($p = 1, 2, \dots, P$) will be combined into a new unified node features matrix $\left(H^{(m)}\right)^{(0)}$ as the initial node features with

$$\left(H^{(m)}\right)^{(0)}_i = \frac{1}{\sum_{p=1}^P I[\mathcal{V}_i \in S_p]} \sum_{p=1}^P \left(H_p^{(m)}\right)_i \cdot I[\mathcal{V}_i \in S_p] \quad (3)$$

, where $\mathcal{V}_i$ represents the node/vertex i in the graph and $I[\mathcal{V}_i \in S_p]$ is the indicator function, whose value will be one if $\mathcal{V}_i \in S_p$. In the end, the initial node features for global level propagation will be $\left(H^{(m)}\right)^{(0)} \in \mathbb{R}^{n \times d^{(0)}}$, $d^{(0)} = d' + d$. Then $\left(H^{(m)}\right)^{(L)}$ will be generated by weighted bi-directional message propagation via

$$(H^{(m)})^{(L)} = \text{MPN} \left((H^{(m)})^{(0)} \right) \quad (4)$$

, where MPN is the global bi-directional message propagation network (see **Appendix B.2** for details).

2.2 Downstream tasks

Drug combination response predictions by decoding drug node embeddings. After obtaining the embedded node features $(H^{(m)})^{(L)} \in \mathbb{R}^{n \times 3d^{(L)}}$ from the global message passing network, the features for drug A are represented as $(H_{drugA}^{(m)})^{(L)} \in \mathbb{R}^{1 \times 3d^{(L)}}$ and the features for drug B are represented as $(H_{drugB}^{(m)})^{(L)} \in \mathbb{R}^{1 \times 3d^{(L)}}$. Utilizing the decagon decoder⁴⁷, the prediction of combo score will be calculated in the following equation:

$$g \left((H_{drugA}^{(m)})^{(L)}, (H_{drugB}^{(m)})^{(L)} \right) = (H_{drugA}^{(m)})^{(L)} D U D^T ((H_{drugB}^{(m)})^{(L)})^T \quad (5)$$

, where $D \in \mathbb{R}^{3d^{(L)} \times E}$ and $U \in \mathbb{R}^{E \times E}$ are trainable decoder matrices, as illustrated in Figure 1 from step 4.

Patient outcome predictions. With the embedded node features, the global mean pooling strategy was applied to predict the patient outcome with

$$\widehat{y^{(m)}} = \arg \max \left(\text{MLP} \left(\text{AVG} \left[(H^{(m)})^{(L)} \right] \right) \right) \quad (6)$$

, where $\widehat{y^{(m)}} \in \mathbb{R}^C$ and C is the number of sample types (see Figure 1 step 4).

2.3 Interpretation using subgraph attention-based score

As aforementioned, we added the K -hop subgraph message propagation with attention in each subgraph, and the attentions or weights on each edge can potentially indicate and interpret the signaling flow on the signaling network to affect the drug combination response. Specifically,

putting the attentions/weights on each edge will use the trained parameters $a \in \mathbb{R}^{2d'}$ and $W \in \mathbb{R}^{d \times d'}$. And for cell lines set, the set $\mathcal{C} = \{C_1, C_2, \dots, C_r, \dots, C_R\}$ contains $R = 39$ cell lines; for ROSMAP AD dataset, $R = 128$, and we can choose the input node feature matrix within a specific cell line or sample to form the cell-line / sample specific attention matrix in k -hop by

$$\left(\alpha_{ij}^{(C_r)}\right)^{(k)} = \sum_{p=1}^P \left(\alpha_{ij}^{(m)}\right)_p^{(k)}, (\forall X^{(m)} \in C_r) \quad (7)$$

, where $(\alpha_{ij}^{(C_r)})^{(k)}$ will be the element in $(\mathcal{A}^{(C_r)})^{(k)}$, the k -hop attention matrix for cell line C_r . And above formula shows the calculation for the 1-head attention matrix. For the multi-head attention mechanism, the averaged value will be used for the attention matrix.

3. Results

3.1 Experimental Setup

To validate the proposed model on drug prediction task, we implemented a 5-fold cross-validation approach. This method involved the use of a four-element vector, $\langle D_A, D_B, C_C, S_{ABC} \rangle$, as the model input. In this tuple, D_A and D_B denote the specific drug combinations being analyzed. C_C represents the cell line name, which is integrated with multi-omic data to provide comprehensive contextual information. Lastly, S_{ABC} reflects the synergistic drug effect observed on the cell line C_C . In total, there are 2788 model input data points for the NCI ALMANAC dataset and 1008 model input data points for the O'Neil dataset. In each fold model, 4 folds of the data points will be used as the training dataset, and the rest 1-fold will be used as the testing dataset. For those 1489 genes, their 8 features indicate the RNA sequence number, copy number variation, gene amplification, gene deletion, gene methylation maximum and minimum value and whether they have a connection spanning D_A and D_B respectively. Regarding the AD samples outcome prediction on ROSMAP dataset, we utilized 138 samples from the ROSMAP dataset, categorized

by disease status (74 AD, 64 non-AD). To address the data imbalance, we performed downsampling for the classification task. For the AD vs. non-AD classification, we downsampled the AD samples to match the 64 non-AD samples, resulting in a dataset with 64 AD and 64 non-AD samples. Afterwards, 5-fold cross-validation was leveraged to evaluate the performance of our model. For those 2099 genes, their 10 features indicate the methylation values on upstream, distal promoters, proximal promoters, core promoters and downstream, the genetic mutations with (number of duplicates, deletions and mCNV), the gene expression values and protein expression values. For each drug pair, their 8 features in drug prediction task or 10 features in ROSMAP sample outcome prediction task were initiated with zeros. Furthermore, to indicate connections between nodes, adjacency matrices were also created. Those matrices were formed from the KEGG, which contains gene pairs with sources and destinations. For drug-gene edges, the connections are bidirectional, which means that in adjacency matrices, those elements are symmetric.

3.2 Hyperparameters

Subsequently, a model was developed by using pytorch and torch geometric. For both drug prediction and ROSMAP sample outcome prediction tasks, the learning rate started at 0.002 and was reduced equally within each batch for a certain epoch stage. And the epochs after 60 will keep the learning rate at 0.0001. Adam optimizer was chosen for optimization with $\text{eps}=1\text{e-}7$ and $\text{weight_decay}=1\text{e-}20$. We empirically set the ***K-hop Subgraph Message Propagation*** part with $K = 3$ and the ***Global Bi-directional Message Propagation*** part with $L = 3$. Afterward, the feature dimensions will vary at the different layers and will be denoted by: $(d^{(1)}, d^{(2)}, \dots, d^{(l)}, \dots, d^{(L)})$. At the global message propagation part, the layer will concatenate biased node features and transformed node features of the previous layer for both upstream and downstream, generating concatenated dimensions for the output dims being $3 \times d^{(l)}$ in l -th layer.

Output dims of the previous layer served as the input dims for the current layer, as follows: (1) First layer (input dims, output dims): $(d^{(0)}, 3d^{(1)})$; (2) Second layer (input dims, output dims): $(3d^{(1)}, 3d^{(2)})$; (3) Third layer (input dims, output dims): $(3d^{(2)}, 3d^{(3)})$. The final embedded drug node dims were $3d^{(3)}$, $(L = 3)$. For drug prediction task, the decoder trainable transformation matrix dims: $D \in \mathbb{R}^{3d^{(L)} \times E}$ and $U \in \mathbb{R}^{E \times E}$ were used as trainable decoder matrices, with changeable parameters of E to adapt model performance and the model used $E = 150$. In ROSMAP sample prediction task, the trainable graph mean pooling was leveraged to predict the sample outcome. As for the LeakyReLU function, the parameter α was set as 0.1 for both tasks.

3.3 M3NetFlow improves drug combination synergy and AD prediction accuracy.

To evaluate the model performance in terms of synergy score prediction for drug combinations and predictions on ROSMAP AD samples, we conducted 5-fold cross-validation. As shown in **Table 3**, the average prediction (using the Pearson correlation coefficient), was about 61% Pearson correlation using the test data in the NCI ALMANAC dataset and was about 64% Pearson correlation using the test data in the O'Neil dataset. Regarding the ROSMAP dataset, the average prediction accuracy was about 66% using the test data in the ROSMAP dataset. These prediction results are comparable with existing deep learning models^{38,39}. Moreover, we also compared our proposed model M3NetFlow with other deep learning models, which included the GCN⁴⁰, Graph Attention network⁴¹ (GAT), UniMP⁴², MixHop³⁰, Principal Neighborhood Aggregation⁴³ (PNA) and GIN⁴⁴. By checking the p values over 5-fold cross validation, the performances of the M3NetFlow have significant improvement over most of the GNN-based methods (see **Table 3** and **Figure 2a**).

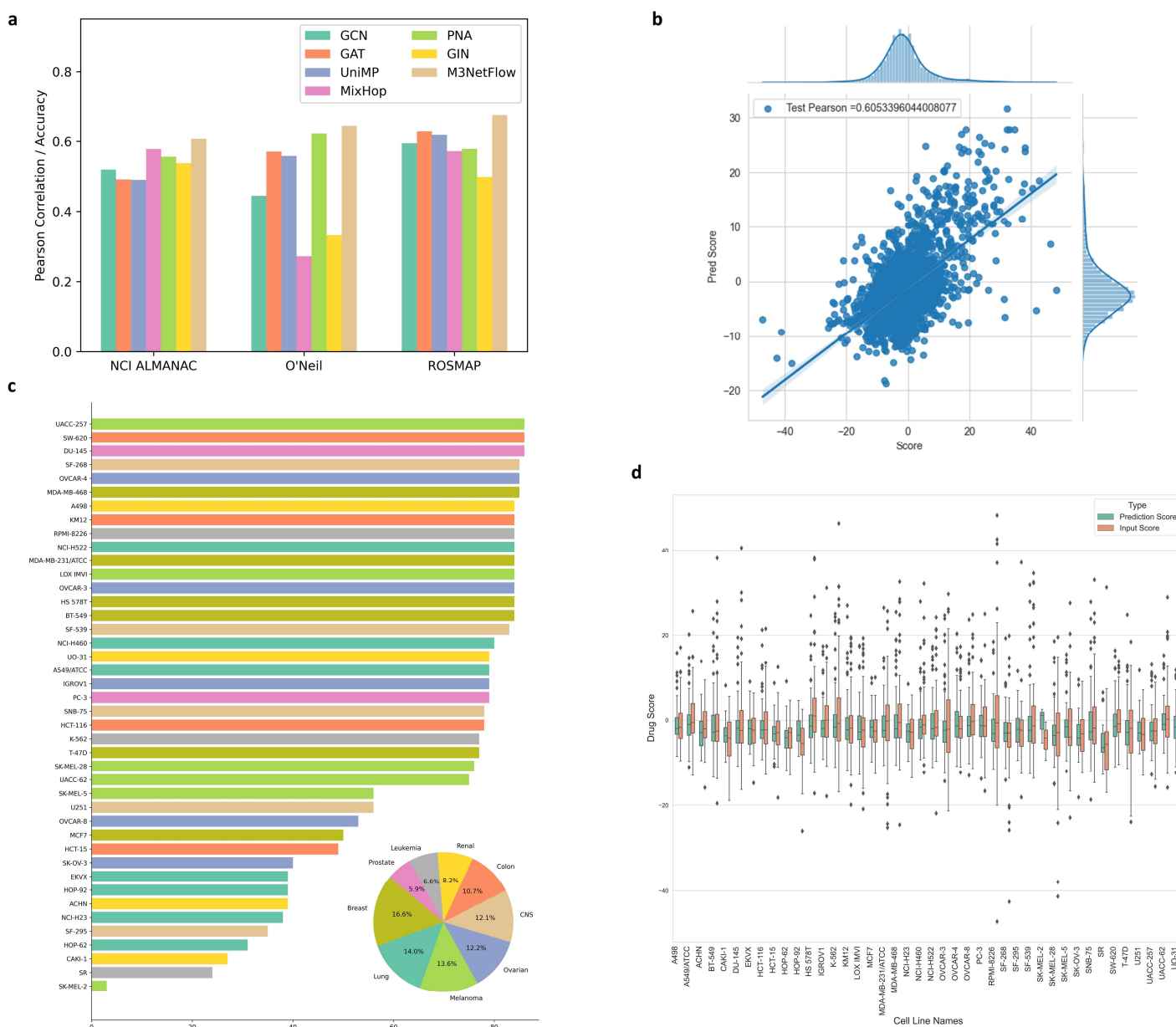


Figure 2. Model performance and overview of input dataset NCI ALMANAC, O'Neil (drug combination multi-omic data) and ROSMAP (AD multi-omic data). **a.** Pearson correlation comparisons for GCN, GAT, UniMP, MixHop, PNA, GIN and **M3NetFlow** models for NCI ALMANAC, O'Neil and ROSMAP dataset. **b.** Scatter plot of the model with data points in whole NCI ALMANAC dataset. **c.** Distributions of all cell lines in whole NCI ALMANAC dataset. **d.** Box plots across all cell lines in the whole NCI ALMANAC dataset.

Table 3. Model comparisons using average Pearson correlation and prediction accuracy of 5-fold cross-validation using NCI ALMANAC, O’Neil datasets and ROSMAP datasets.

Dataset	NCI ALMANAC	NCI ALMANAC	O’Neil	O’Neil	ROSMAP	ROSMAP
	P values		P values		P values	
GCN	51.93% ± 3.24%	0.006670411	44.47% ± 6.02%	0.002229	59.43% ± 4.53%	0.049650832
GAT	49.16% ± 2.07%	0.000221277	57.06% ± 3.40%	0.014098	62.80% ± 7.11%	0.332946723
UniMP	49.02% ± 4.21%	0.006497083	55.84% ± 9.98%	0.196887	61.83% ± 3.78%	0.124477379
MixHop	57.78% ± 3.34%	0.185206079	27.15% ± 9.14%	0.00123	57.20% ± 3.92%	0.014849014
PNA	55.63% ± 2.26%	0.01219272	62.20% ± 2.07%	0.240673	57.83% ± 1.82%	0.017236825
GIN	53.76% ± 2.25%	0.003440775	33.12% ± 9.03%	0.002405	49.83% ± 5.71%	0.001882831
M3NetFlow	60.72% ± 0.70%	-	64.36% ± 2.31%	-	67.34% ± 5.12%	-

3.4 M3NetFlow ranks important targets via attention score

Based on the attention matrix of each sample (like a cell line or an AD patient sample) between neighboring nodes (genes) on the graph in each fold of cross validation, the average attention matrix in was calculated. Therefore, the weighted importance of each node (gene) of each sample will be calculated based on the attention with

$$\mathcal{D}_g^{(C_r)} = \sum_i^n \bar{\mathcal{A}}_{ig}^{(C_r)} + \sum_j^n \bar{\mathcal{A}}_{gj}^{(C_r)} \quad (8)$$

, where $\bar{\mathcal{A}}_{ig}^{(C_r)}$ is the element of averaged 5-fold attention matrix in 1st hop from a sample C_r in the i -th row and j -th column and $\mathcal{D}_g^{(C_r)}$ is the node importance score for gene g in the sample C_r . Combining all the node importance vectors, the matrix $\mathcal{D} \in \mathbb{R}^{n \times R}$ will be generated, where n is the number of nodes, and R is the number of samples. However, some of genes may show the importance of relatively higher scores in each sample, which weakens the analysis of the sample-specific analysis. In this way, the idea of reweighting the gene importance score was created based on the TF-IDF to reweight the node importance in each sample with

$$\mathcal{W}_g = \text{Log} \left(\frac{R + 1}{S_g + 0.01} \right) \quad (9)$$

$$\mathcal{D}'_g = \mathcal{W}_g \cdot \mathcal{D}_g \quad (10)$$

, where S_g is the number of samples with a higher node importance score than the threshold S in all samples. Here, the 95 percentile of node degree values in the whole matrix $\mathcal{D}$ was used for S . Finally, the reweighted node degree (importance score) matrix $\mathcal{D}'$ will be generated. Based on the reweighted matrix $\mathcal{D}'$, the node (gene) importance score will be calculated. The sum of importance score of individual genes across samples of disease will be used as the disease importance score of genes.

3.5 Case 1: Synergistic drug combinations are associated with the identified key targets

Based on the attention score calculated for each specific cell line in the section 3.4, the highest 5, and lowest 5 drug scores and their corresponding drugs and drug-targeted genes were collected for each cancer cell lines in NCI ALMANAC dataset (check **Figure 3a**). Comparing the targeted genes by drugs from the highest 5 drug scores and lowest drug scores, the differences are statistically significant with p values in more than half cell lines. Since the SK-MEL-2 cell line only has 3 drug screen data points and all of them are smaller than 0, this cell line was excluded from the statistical analysis. In total, there are 27 out of 41 (~65%) cell lines have p values smaller than 0.1 (filled with deep orange color in **Table S1**) and 32 out of 41 (~78%) cell lines have p value smaller than 0.3 (filled with light orange color in **Table S1**). And 37 out of 41 (~90%) cell lines have higher gene importance scores on both mean and median values (filled with light green in **Table S1**). **Figures 4-6** provided more explicit evidence of a trend where genes targeted by drugs with higher scores exhibited higher node degrees (gene importance scores). In each cell line, the left boxplots compared the node degree distribution of genes targeted by the top 5 drugs with the highest scores against those targeted by the bottom 5 drugs with the lowest scores. The right boxplots in each cell line represented the overall comparisons between the top 5 and bottom

5 drug scores. Consequently, we could assign prior targets to genes with the highest node degree (gene importance scores) using our calculations. This allowed us to rank the genes in each cell line according to their node degree (importance scores), which can be found in **Table S2**. We generated a list of the top 20 genes in each cell line, available in **Table S3**. Subsequently, we identified overlapping genes among the cell lines of the same cancer type based on the top 20 genes in each line for cancer-specific analysis, as depicted in **Figure S1** from Appendix Section C.

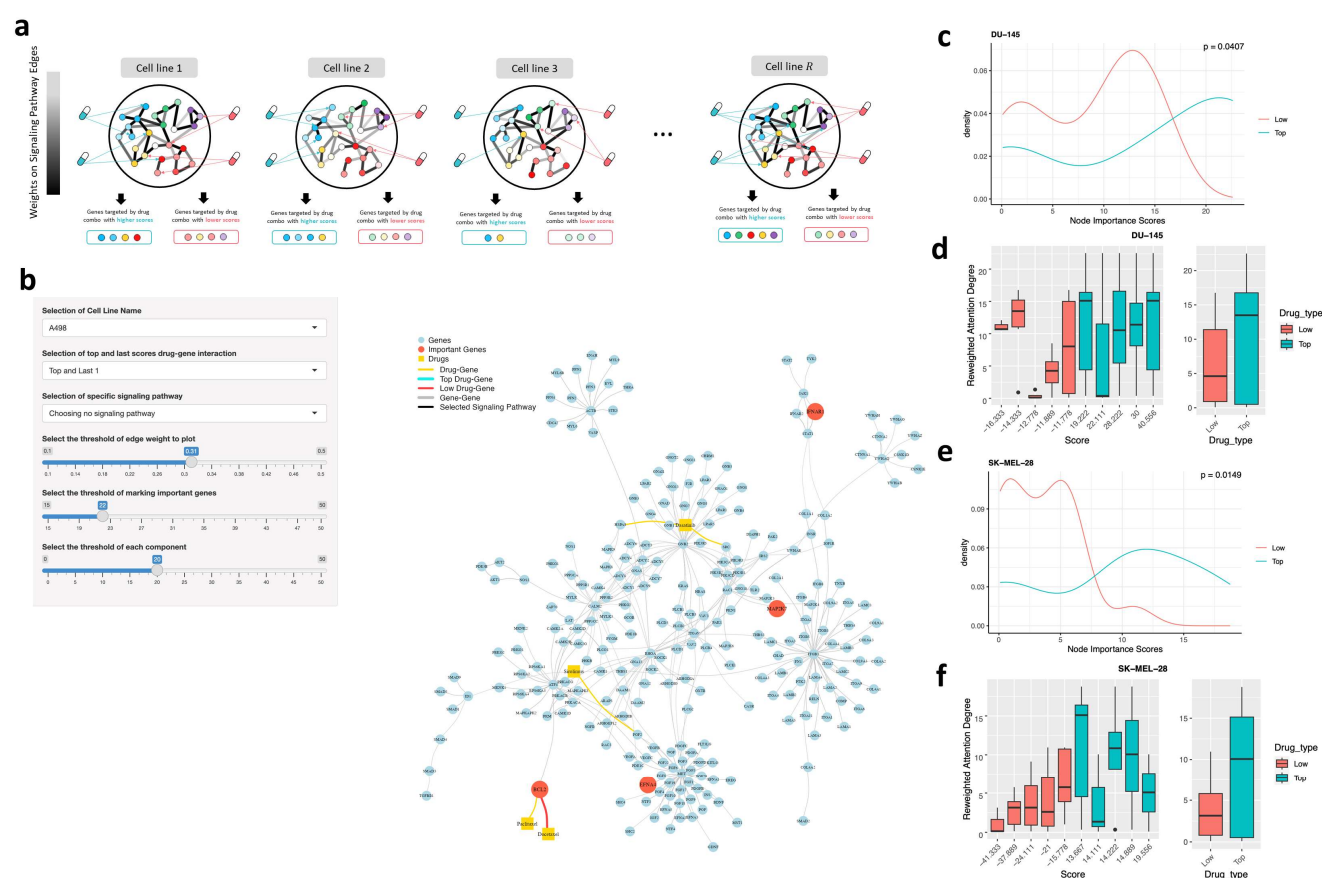


Figure 3. Validation and visualization through important node analysis **a.** Procedures of analyzing the important genes targeted by the highest 5 and lowest 5 drug combinations **b.** Visualization tool **NetFlowVis** (check Appendix Section D for details of this tool) for core signaling network interactions of cell line DU-145 **c-f.** Density and box plots by comparing the genes targeted by the highest 5 and lowest 5 drug combinations in DU-145 and SK-MEL-28 cell lines.

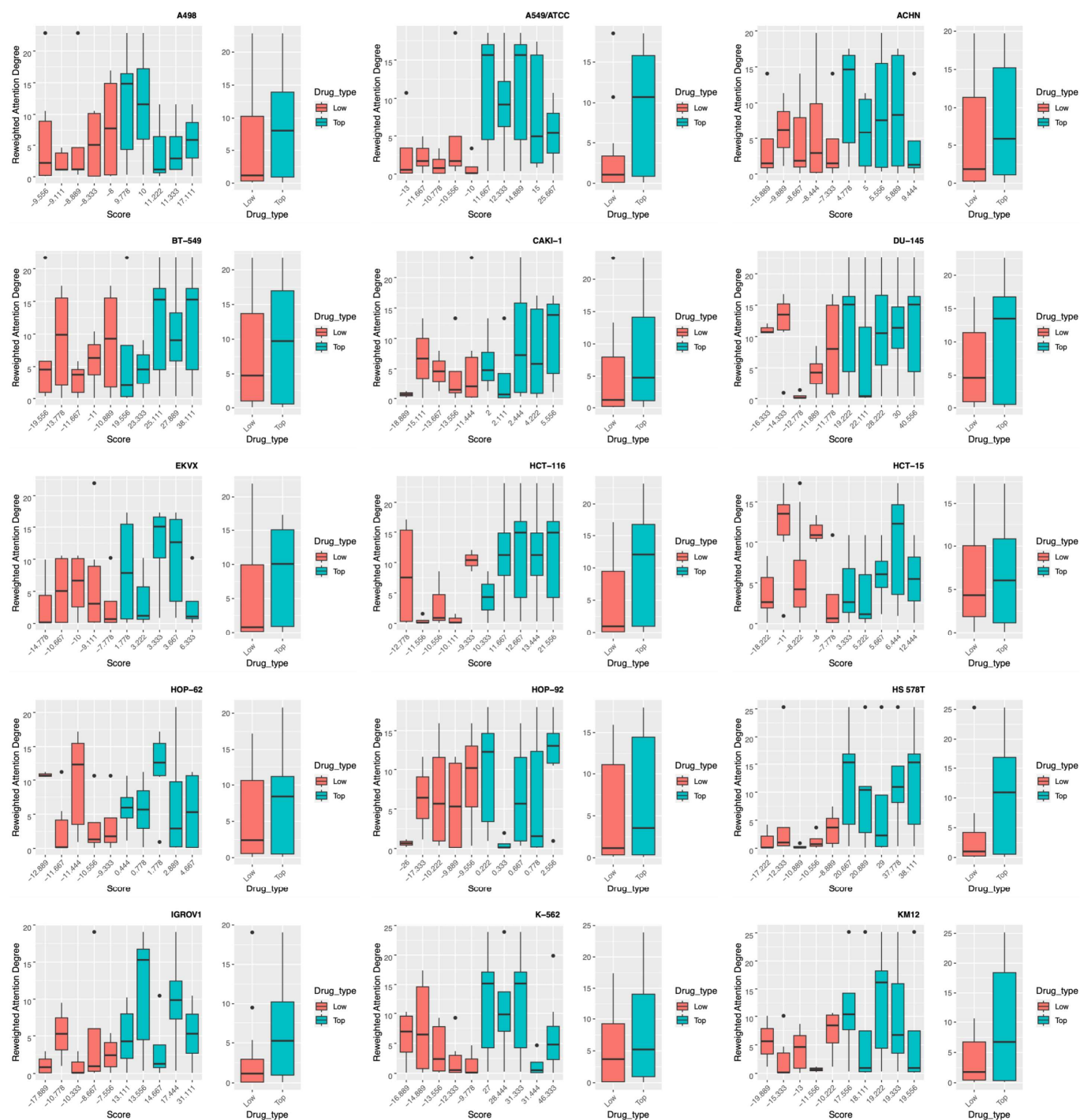


Figure 4. Cell line A498, A549/ATCC, ACHN, BT-549, CAKI-1, DU-145, EKVX, HCT-116, HCT-15, HOP-62, HOP-92, HS 578T, IGROV1, K-562, KM12 genes degree distribution targeted by highest 5 and lowest 5 drug scores.

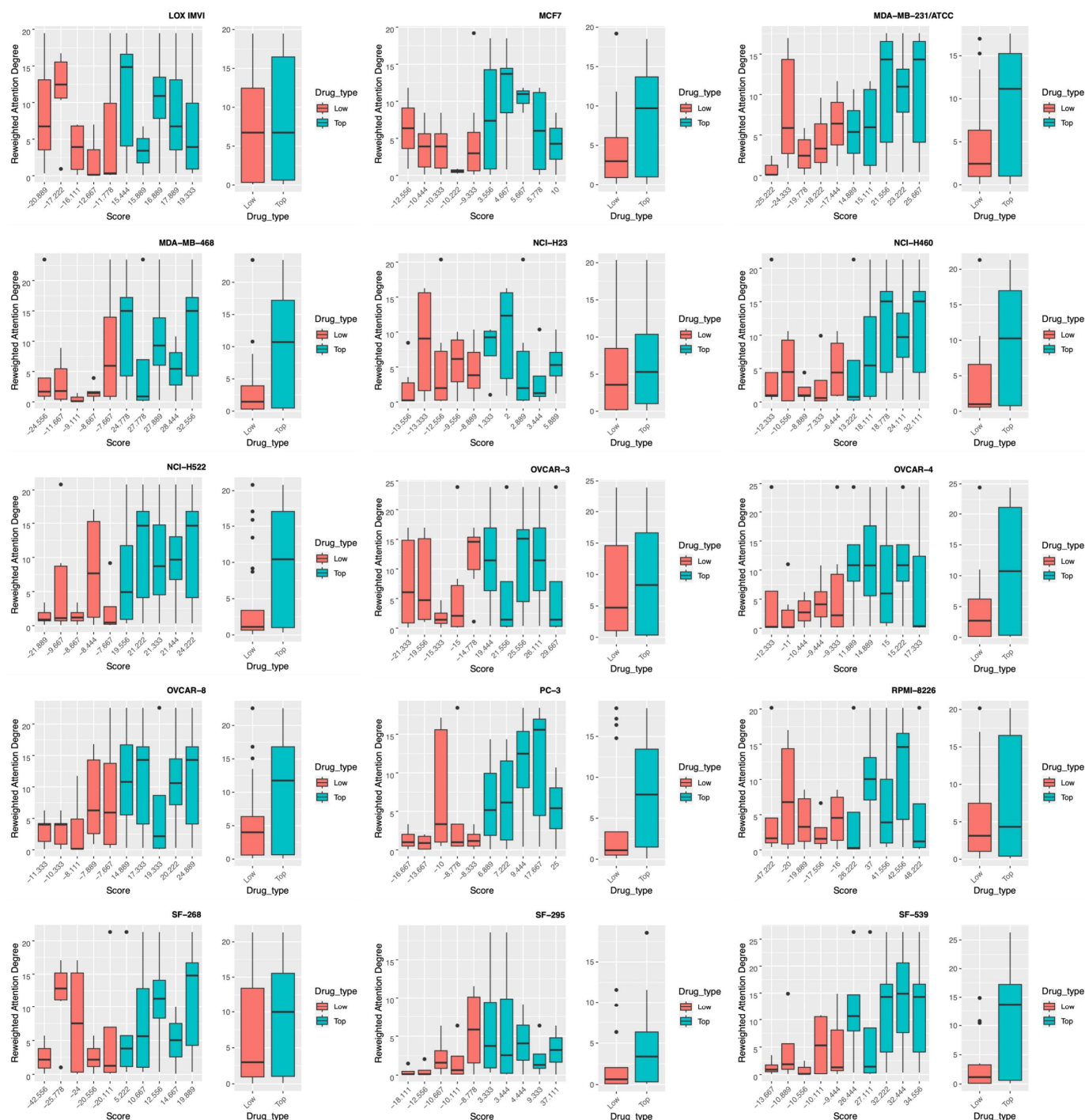


Figure 5. Cell line LOX IMVI, MCF7, MDA-MB-231/ATCC, MDA-MB-468, NCI-H23, HCI-H460, NCI-H522, OVCAR-3, OVCAR-4, OVCAR-8, PC-3, RPMI-8226, SF-268, SF-295, SF-539 genes degree distribution targeted by highest 5 and lowest 5 drug scores.

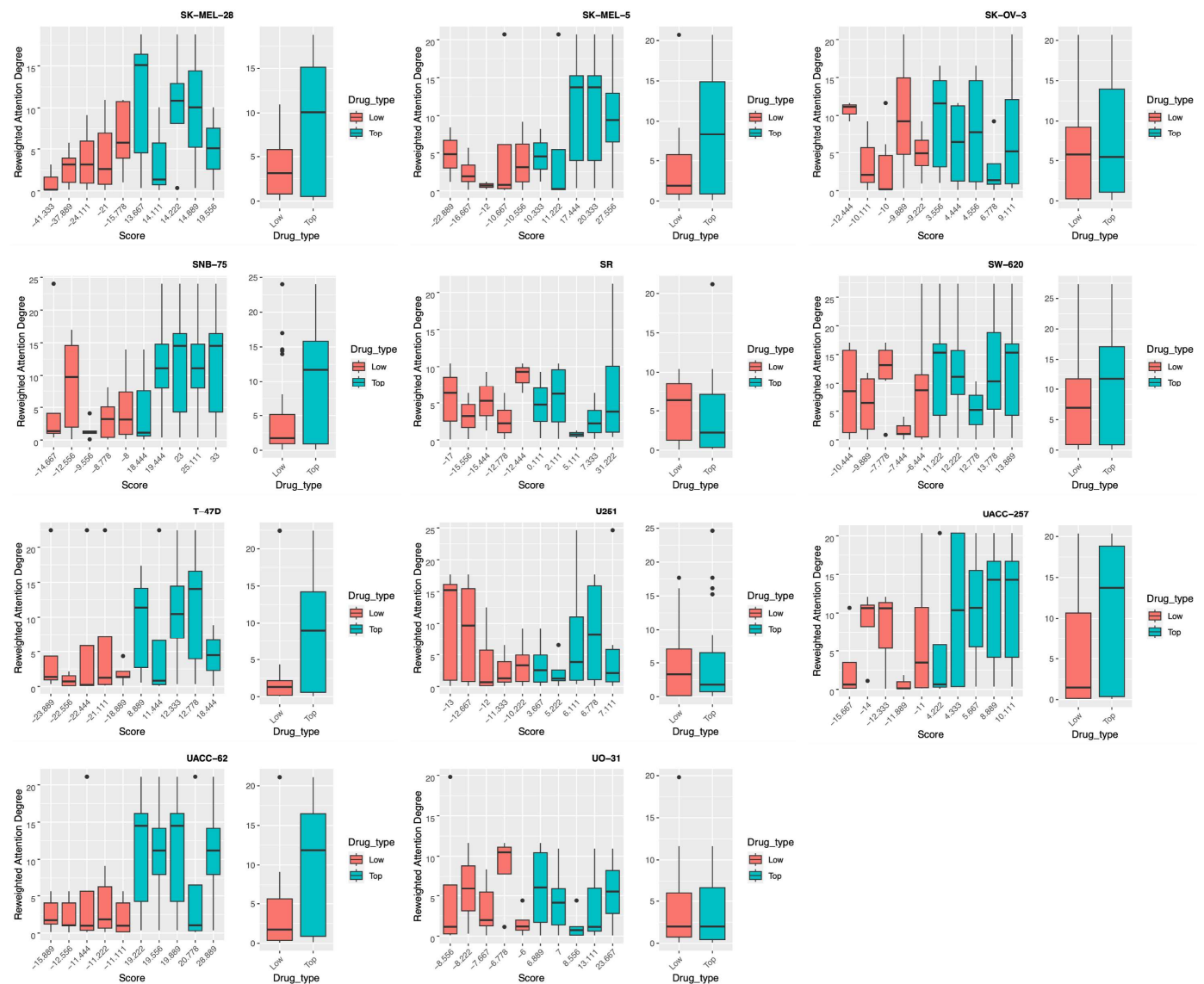


Figure 6. Cell line SK-MEL-28, SK-MEL-5, SK-OV-3, SNB-75, SR, SW-620, T-47D, U251, UACC-257, UACC-62, UO-31 genes degree distribution targeted by highest 5 and lowest 5 drug scores.

3.6 Case 2: Alzheimer's disease associated biomarkers and pathways

Generating attention-based scores from **M3NetFlow** (see Section 3.4), the sample specific edge weight matrices will be aggregated by

$$\bar{\mathcal{A}}^{(AD)} = \frac{1}{R} \sum_{r=1}^R \bar{\mathcal{A}}^{(C_r)} \quad (11)$$

, where $R = 64$ is the number of AD samples in the set $\mathcal{C}$ and $\bar{\mathcal{A}}^{(AD)} \in \mathbb{R}^{n \times n}$ is the aggregated

edge weight for AD patients. By setting the filters (edge threshold as 0.106 and a small component threshold as 15), we identified 100 potential important genes for AD. Among those genes, 28 genes are filtered out by setting the threshold of attention-based node weight as 2.0, and 15 of them with p values smaller than 0.1 in at least one of the 10 multi-omic features (see **Figure 7a-b**). To evaluate the top targets ranked by attention-based node weight, the top-ranked 28 AD associated biomarkers were further analyzed via pathway enrichment analysis. Interestingly, a set of AD associated signaling pathways are identified. As shown in **Figure 7c**, the top-ranked targets are involved in a set of signaling transduction pathways, which indicates the importance of these targets.

Among them, several signaling pathways play critical roles in Alzheimer's disease (AD), particularly through mechanisms involving inflammation, immune responses, and cellular growth and death. For instances, the B cell receptor (BCR) and T cell receptor (TCR) signaling pathways are vital for B cell and T cell activation, which plays a crucial role in immune surveillance and inflammation. Shared molecular components between the BCR and TCR pathways, including MAP2K1, JUN, CHUK, RELA, NFKB1, IKBKB, and AKT genes, suggest significant cross-talk that contributes to neuroinflammatory processes in AD. Prolonged activation of these immune pathways may exacerbate chronic inflammation and contribute to AD pathology by sustaining harmful neuroinflammatory responses⁴⁵. The NF- κ B signaling axis, a critical component in both BCR and TCR pathways, is particularly relevant in AD due to its role in regulating inflammatory cytokine production. Chronic activation of NF- κ B has been linked to increased amyloid-beta (A β) deposition and neurofibrillary tangle formation, both of which are hallmark features of AD pathology^{46,47}. Additionally, activation of NF- κ B, MAPK, and AKT signaling cascades can induce neuronal apoptosis, leading to cognitive decline associated with the disease.

The MAPK, RAS, PI3K-Akt, and FoxO signaling pathways also play critical roles in modulating

immune responses and neuronal survival. The p38 MAPK pathway, for example, drives neuroinflammation by activating microglia and promoting the release of pro-inflammatory cytokines, which can result in neuronal death⁴⁸. Dysregulation of the RAS-RAF-MEK-ERK signaling cascade, another crucial pathway in AD, affects neuronal survival and apoptosis. Overactivation of RAS signaling increases oxidative stress, contributing to A β accumulation and tau pathology, both of which are central to neurodegeneration in AD⁴⁹. While A β accumulation is an early event in AD, tau pathology is more closely associated with cognitive decline, and together, these processes are considered the primary drivers of neuronal death in AD^{50,51}.

The Akt signaling pathway, which promotes cell survival by inhibiting apoptosis through downstream effectors such as BCL2, is also impaired in AD. Reduced Akt activity has been linked to increased neuronal apoptosis and the accumulation of A β and hyperphosphorylated tau. Notably, the PI3K-Akt pathway is intimately connected to insulin signaling, which is disrupted in AD and is characterized by reduced Akt signaling, impaired glucose metabolism, and an increased vulnerability to neurodegeneration⁵². Furthermore, dysregulation of the FoxO transcription factors due to impaired PI3K-Akt signaling leads to enhanced expression of pro-apoptotic genes such as BIM and PUMA, contributing to synaptic and neuronal loss⁵³.

Neurotrophin signaling, which regulates axonal growth and regeneration via MAPK and PI3K-Akt pathways, is another pathway that becomes disrupted in AD. This disruption impairs axonal repair mechanisms and contributes to synaptic loss and neuronal degeneration⁵⁴. Additionally, endocrine-related pathways, including insulin and relaxin signaling, play crucial roles in maintaining cellular homeostasis and survival. Dysregulation of these pathways in AD contributes to neuronal apoptosis, neuroinflammation, and impaired protein clearance, further exacerbating disease pathology^{52,55,56}.

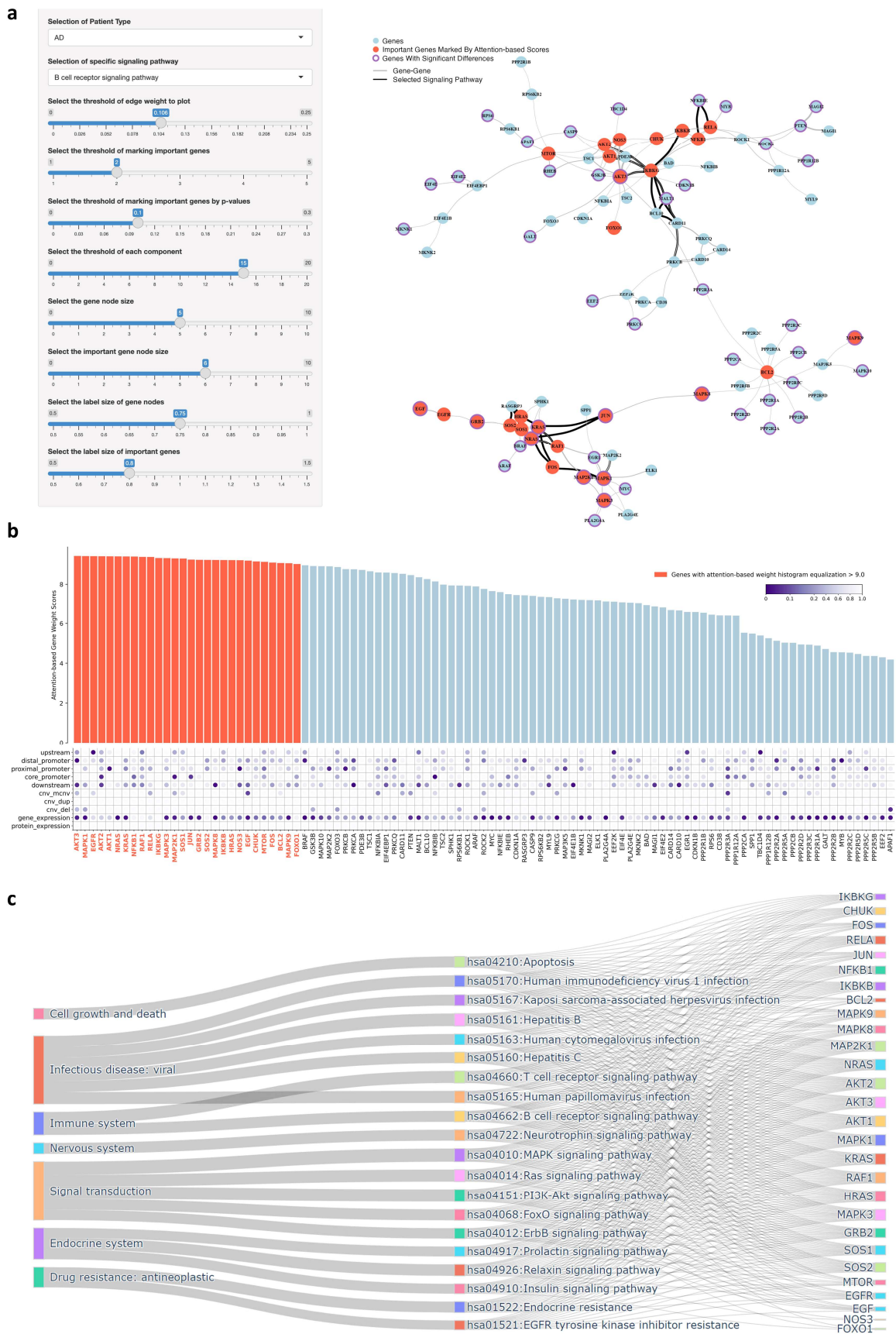


Figure 7. Biomarkers and pathways associated with Alzheimer's disease. **a.** Visualization tool **NetFlowVis** for core signaling network interactions of AD. The important nodes are set as red filtered by attention-based node weight threshold of 2.0 and nodes with significant (p values <0.1) differences between AD and non-AD samples are circled with purples **b.** Barplot of the attention-based importance score measured by **M3NetFlow**. The gene names and bars marked with red are important genes selected by setting attention-based node weight as 2.0, and the borders of the bar are set as purple if significant (p values <0.1) differences existed between AD and non-AD samples. **c.** Sankey plot of enriched signaling pathways of the identified AD associated genes selected by attention-based node weight (threshold as 2.0).

4. Discussion and conclusion

Along with the advancement of next-generation sequencing (NGS) technology, multi-omic data of diseases are being generated to characterize the dysfunctional molecular targets and signaling pathways of complex diseases, like cancer and AD, which are valuable and essential for the development of personalized medicine or precision medicine prediction. However, it remains a challenging task to identify key molecular targets and core signaling pathways from a large signaling network with thousands of signaling targets and extensive signaling interactions. Herein, we present a novel graph AI model, named M3NetFlow, for generic multi-omic data integrative analysis. In this model, multi-omic data are used as the numerical features of individual proteins, and the protein-protein interactions form a large-scale signaling graph. Unlike existing graph models, we divided the large-scale network into signaling modules and applied a multi-hop strategy to better infer the essential proteins and their interactions. The experimental evaluation on two real multi-omic data analysis applications showed that the proposed model outperformed existing graph models in accuracy and is able to identify essential targets and core signaling pathways, demonstrating its interpretability for Alzheimer's disease or drug combination synergy. The proposed model can be applied to 1) anchor-target guided or 2) generic target and pathway inference, as indicated by the two multi-omic data analysis applications that represent widely

conducted tasks in omic data analysis. In the first scenario, a set of targets of interest, such as drug targets or known molecular targets, are used as the anchor nodes. Then, the signaling flows/information from neighboring graph nodes propagate to the anchor nodes. The embeddings of the anchor nodes are used as input for decoders to predict sample classes or drug combination responses. In the second scenario, the embeddings of graph nodes are pooled together as input for decoders to predict sample classes. Based on the attention scores of graph nodes, the essential targets and signaling pathways are further identified. Therefore, the proposed model can be applied for generic, integrative, and interpretable multi-omic data analysis tasks, and the code is publicly accessible.

It is still an exploratory study for multi-omic data analysis. There are some limitations that need further investigation. For example, more signaling pathways and larger protein-protein interactions should be evaluated. Moreover, dividing large signaling graphs into subnetworks or network modules can be achieved by using biologically meaningful annotations, such as gene ontology (GO) terms. Additionally, more multi-omic datasets are being generated. Combining multi-omic data from different diseases can provide a larger sample size than individual disease datasets, which could improve the training or pre-training of graph AI models and help identify pan-disease or disease-specific targets. It is also interesting to expand graph models from tissue-level multi-omic data to single-cell multi-omic data, which can be more challenging due to the large number of single-cell samples. Therefore, novel and improved graph AI models are needed to integrate and interpret multi-omic datasets, identify and infer key molecular targets and signaling pathways of complex diseases, and guide the development of precision medicine.

Appendix

Section A. Data collection and preprocessing

A.1 Drug combination screening datasets

In this paper, we utilized the NCI ALMANAC and O'Neil⁵⁷ drug screening datasets to train the deep learning models. The NCI ALMANAC dataset consists of combo scores for various permutations of 104 FDA-approved drugs, representing their impact on tumor growth in NCI60 human tumor cell lines. For evaluating the synergy score of two drugs on a specific tumor cell line, we utilized the average combo-score of two drugs at different doses, employing a 4-element tuple: $\langle D_A, D_B, C_C, S_{ABC} \rangle$. On the other hand, the O'Neil dataset, obtained from the DrugComb⁵⁸ platform, served as another drug combination screening dataset. We extracted the processed dataset from this platform, which provided average synergy Loewe scores. These scores assessed the synergy score of two drugs on a given tumor cell line, employing a 4-element tuple: $\langle D_A, D_B, C_C, S_{ABC} \rangle$.

A.2 Multi-omic data of cancer cell lines

In this study, multi-omic data comprising RNA-seq, copy number variation, gene methylation, and gene mutation data for a total of 1,489 genes were incorporated into the model. These data were obtained from the Cell Model Passports⁵⁹ and CCLE⁶⁰ databases. Through the identification of overlapping cell lines from these databases, we identified 42 cell lines that were present in both the Cell Model Passports database and the NCI ALMANAC dataset, which are A498, A549/ATCC, ACHN, BT-549, CAKI-1, DU-145, EKVX, HCT-116, HCT-15, HOP-62, HOP-92, HS 578T, IGROV1, K-562, KM12, LOX IMVI, MCF7, MDA-MB-231/ATCC, MDA-MB-468, NCI-H23, NCI-H460, NCI-H522, OVCAR-3, OVCAR-4, OVCAR-8, PC-3, RPMI-8226, SF-268, SF-295, SF-539, SK-MEL-2, SK-MEL-28, SK-MEL-5, SK-OV-3, SNB-75, SR, SW-620, T-47D, U251, UACC-257, UACC-62, UO-31. Additionally, we found 24 cell lines that were present in both the Cell Model Passports database and the O'Neil dataset, which are A2058, A2780, A375, CAO3V, HCT116, HT144, LOVO, MDAMB436, NCI-H460, NCIH1650, NCIH2122, NCIH23, OV90, OVCAR3, RKO, RPMI7951, SK-OV-3, SKMEL30, SKMES1, SW-620, SW837, T-47D, UACC62, VCAP.

A.3 Multi-omic data and clinical phenotypes datasets from ROSMAP

Following the acquisition of the datasets, they were reformatted into 2-dimensional data frames, with columns dedicated to sample identifiers, such as IDs and names, and rows corresponding to probes, gene symbols, and gene IDs. To successfully integrate the multi-omic data with clinical information, it was essential to match identical samples across the various datasets. This required standardizing the row data—including probes, gene symbols, and gene IDs—into a unified gene-level format, either by aggregating gene-specific measurements or by resolving duplicates from gene synonyms. Genes were subsequently mapped to a reference genome to ensure precise annotation within the multi-omic datasets. Gene counts were then normalized across the datasets, with missing values imputed using zeros or negative ones as needed. After aligning all columns to standard sample IDs and rows to standardized gene IDs, and ensuring a consistent number of samples and genes, the data was prepared for integration into Graph Neural Network (GNN) models, where epigenomic, genomic, transcriptomic and proteomic data were employed as features for nodes.

A.4 KEGG Signaling Pathways

About 59,241 gene-gene interactions for over 8,000 genes across different signaling pathways were collected from KEGG⁶¹ database. And 48 signaling pathways in the KEGG dataset were selected as the ground truth of subgraphs in the gene-gene interactions, which were AGE-RAGE signaling pathway in diabetic complications, AMPK signaling pathway, Adipocytokine signaling pathway, Apelin signaling pathway, B cell receptor signaling pathway, C-type lectin receptor signaling pathway, Calcium signaling pathway, Chemokine signaling pathway, ErbB signaling pathway, Estrogen signaling pathway, Fc epsilon RI signaling pathway, FoxO signaling pathway, Glucagon signaling pathway, GnRH signaling pathway, HIF-1 signaling pathway, Hedgehog

signaling pathway, Hippo signaling pathway, Hippo signaling pathway - multiple species, IL-17 signaling pathway, Insulin signaling pathway, JAK-STAT signaling pathway, MAPK signaling pathway, NF-kappa B signaling pathway, NOD-like receptor signaling pathway, Neurotrophin signaling pathway, Notch signaling pathway, Oxytocin signaling pathway, PI3K-Akt signaling pathway, PPAR signaling pathway, Phospholipase D signaling pathway, Prolactin signaling pathway, RIG-I-like receptor signaling pathway, Rap1 signaling pathway, Ras signaling pathway, Relaxin signaling pathway, Signaling pathways regulating pluripotency of stem cells, Sphingolipid signaling pathway, T cell receptor signaling pathway, TGF-beta signaling pathway, TNF signaling pathway, Thyroid hormone signaling pathway, Toll-like receptor signaling pathway, VEGF signaling pathway, Wnt signaling pathway, cAMP signaling pathway, cGMP-PKG signaling pathway, mTOR signaling pathway, p53 signaling pathway. After preprocessing the dataset, 17,259 gene-gene interactions for 1,489 genes were obtained. And 28,843 gene-gene interactions for 2,144 genes were obtained for ROSMAP AD dataset.

A.5 Drug-Target interactions derived from DrugBank

Drug-target information was extracted from the DrugBank⁶² database (version 5.1.5, released 2020-01-03). In total, 15,263 drug-target interactions were obtained for 5435 drugs/investigational agents and 2775 targets. Further, 17 drugs with known targets to the 1489 genes previously identified were selected for use in our model for NCI ALMANAC⁶³, specifically: Celecoxib, Cladribine, Dasatinib, Docetaxel, Everolimus, Fulvestrant, Gefitinib, Lenalidomide, Megestrol acetate, Mitotane, Nilotinib, Paclitaxel, Romidepsin, Sirolimus, Thalidomide, Tretinoin, Vorinostat. And 10 drugs with known targets to the 1489 genes previously identified were selected for use in our model, specifically: Dasatinib, Erlotinib, Lapatinib, Sorafenib, Sunitinib, Vorinostat, Geldanamycin, Metformin, Paclitaxel, Vinblastine.

Section B. Method Details

B.1 K-hop attention-based graph neural network

To incorporate the long-distance information, for each subgraph K -hop attention based mechanism was constructed with

$$\left(\alpha_{ij}^{(m)}\right)_p^{(k)} = \frac{\exp(\text{LeakyReLU}(a[W(X_p^{(m)})_i || W(X_p^{(m)})_j]))}{\sum_{q \in \mathcal{N}_i^{(k)}} \exp(\text{LeakyReLU}(a[W(X_p^{(m)})_i || W(X_p^{(m)})_q]))} \quad (10)$$

$$(H_p^{(m)})_i = \frac{1}{KQ_i} \sum_k^K \sum_{q \in \mathcal{N}_i^{(k)}} [(\alpha_{iq}^{(m)})_p^{(k)} W(X_p^{(m)})_q] \quad (11)$$

, where $(\alpha_{ij}^{(m)})_p^{(k)}$ is the attention score mentioned in equation (1) and $\mathcal{N}_i^{(k)}$ is the neighbor nodes calculated by the adjacency matrix in k -th hop $A^{(k)}$ for the node i (See **Appendix B.3** for the algorithm of k -th hop adjacency matrix). The linear transformation vector $a \in \mathbb{R}^{2d'}$ was also defined. At the same time, the linear transformation for features of each node will be defined as $W \in \mathbb{R}^{d \times d'}$. And $(X_p^{(m)})_t, (H_p^{(m)})_t$ represent the feature and updated feature of node t ($t = 1, 2, \dots, n$). Aside from that, Q_t represents the number of signaling pathway the node t ($t = 1, 2, \dots, n$) belongs to. And above formula shows the calculation for the 1-head attention. The number of head h will be modified in the model, and the node embeddings for each subgraph will take the average of embeddings in every head attention.

B.2 Weighted Bi-directional Message Propagation

The weight matrices are represented as $W_{up}^{(l)} \in \mathbb{R}^{n \times n}$, $W_{down}^{(l)} \in \mathbb{R}^{n \times n}$ for each of the L global message propagation layers, ($l = 1, 2, \dots, L$). The weighted adjacency matrices are defined as follows:

$$A'_{up}^{(l)} = W_{up}^{(l)} \cdot A \quad (12)$$

$$A'_{down}^{(l)} = W_{down}^{(l)} \cdot A^T \quad (13)$$

Here, $A'_{up}^{(l)} \in \mathbb{R}^{n \times n}$ and $A'_{down}^{(l)} \in \mathbb{R}^{n \times n}$ maintain the parameters for drug-gene relationships as in the original adjacency matrix. The mean aggregation of each node's neighbors' features from upstream to downstream from layer $l - 1$ to layer l is achieved using the following equation:

$$\left(H_{up_nei}^{(m)} \right)^{(l)} = D_{in}^{-1} (A'_{up}^{(l)})^T \left(H^{(m)} \right)^{(l-1)} Z_{up}^{(l)} \quad (14)$$

Similarly, the mean aggregation of each node's neighbors' features from downstream to upstream is computed as:

$$\left(H_{down_neigh}^{(m)} \right)^{(l)} = D_{out}^{-1} (A'_{down}^{(l)})^T \left(H^{(m)} \right)^{(l-1)} Z_{down}^{(l)} \quad (15)$$

The biased transformation of nodes and their features is given by:

$$\left(H_{self}^{(m)} \right)^{(l)} = \left(H^{(m)} \right)^{(l-1)} B^{(l)} \quad (16)$$

In these equations, $Z_{up}^{(l)} \in \mathbb{R}^{d^{(l-1)} \times d^{(l)}}$ and $Z_{down}^{(l)} \in \mathbb{R}^{d^{(l-1)} \times d^{(l)}}$ are the linear transformation matrices for each layer l , ($l = 1, 2, \dots, L$). For each layer, the model employs a normalization function $N: \mathbb{R}^{n \times 3d^{(l)}} \rightarrow \mathbb{R}^{n \times 3d^{(l)}}$ and a LeakyReLU activation function with parameter α to map the concatenated node features $\left(H_{concat}^{(m)} \right)^{(l)} = \left(\left(H_{up_nei}^{(m)} \right)^{(l)} \parallel \left(H_{down_neigh}^{(m)} \right)^{(l)} \parallel \left(H_{self}^{(m)} \right)^{(l)} \right) \in \mathbb{R}^{n \times 3d^{(l)}}$ to $H^{(l)} \in \mathbb{R}^{n \times 3d^{(l)}}$ as follows

$$\left(H^{(m)} \right)^{(l)} = \text{LeakyReLU} \left(N \left(\left(H_{concat}^{(m)} \right)^{(l)} \right) \right) \quad (17)$$

, where the normalization function N performs L2 normalization on the demo matrix $V \in \mathbb{R}^{p \times q}$ along the row axis using the equation:

$$v'_{ij} = \frac{v_{ij}}{\sqrt{\sum_{j=1}^q v_{ij}^2}} \quad (18)$$

, where v'_{ij} is an element of the new matrix $V' \in \mathbb{R}^{p \times q}$.

B.3 K-hop adjacency matrix generation algorithm

Algorithm: k -th hop adjacency matrix

Input: initial adjacency matrix $A \in \mathbb{R}^{n \times n}$

Step 1: $A_{sum} = A$

Step 2: For k in $2, 3, \dots, K$:

Step 3: $A^{(k)} = A^k$

Step 4: For i in $1, 2, 3, \dots, n$:

Step 5: For j in $1, 2, 3, \dots, n$:

Step 6: If $A_{ij}^{(k)} > 1$:

Step 7: $A_{ij}^{(k)} = 1$

Step 8: End If

Step 9: If $i = j$:

Step 10: $A_{ij}^{(k)} = A_{ij}^{(k)} - (A_{sum})_{ij} - 1$

Step 11: Else:

Step 12: $A_{ij}^{(k)} = A_{ij}^{(k)} - (A_{sum})_{ij}$

Step 13: End if

Step 14: If $A_{ij}^{(k)} < 0$:

Step 15: $A_{ij}^{(k)} = 0$

Step 16: End If

Step 17: End For

Step 18: End For

Step 19: $A_{sum} = A_{sum} + A^{(k)}$

Step 20: End For

Output: $A^{(2)}, A^{(3)}, \dots, A^{(K)}$

Section C. Downstream analysis of important targets for cancer cell lines

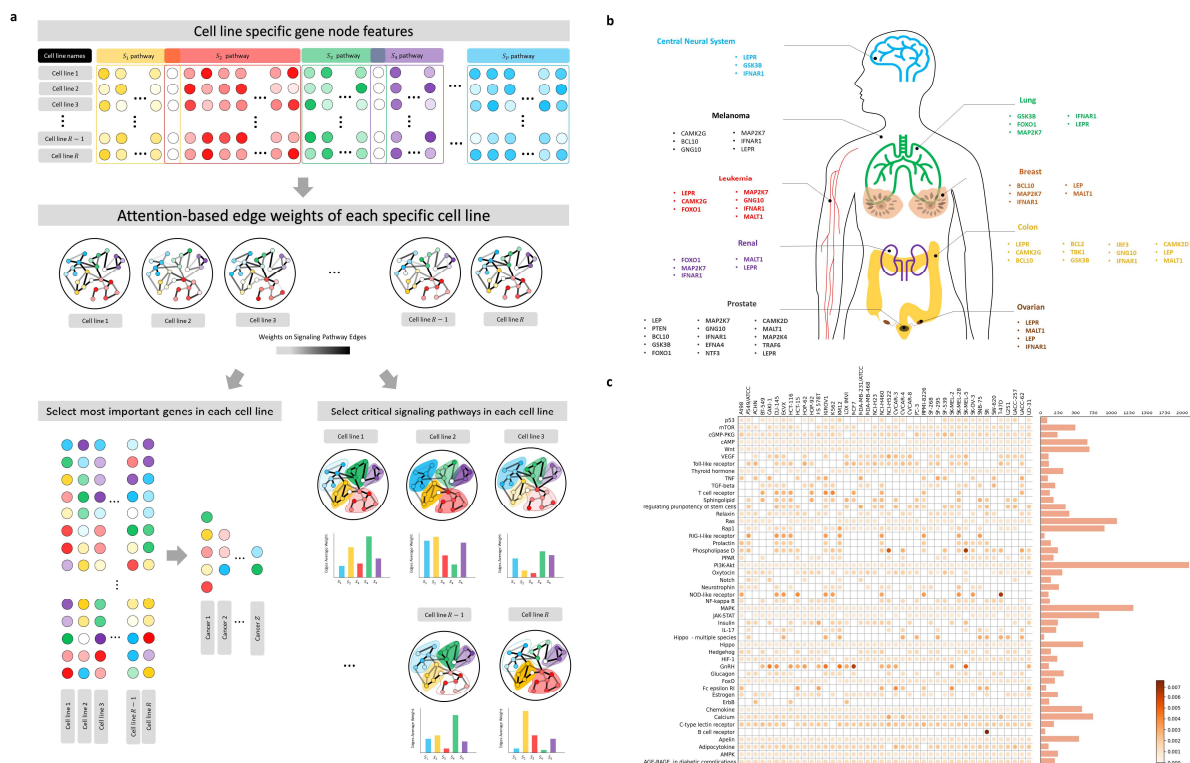


Figure S1. Downstream analysis of important genes and signaling pathways. a. Procedures of identifying the important genes in different cell lines and cancers, and identifying the important signaling pathways in different cell lines. b. Overlapped genes for the cell lines in each type of cancer. c. Heatmap of sum of edge weights in different signaling pathways in whole NCI ALMANAC dataset.

Section D. *NetFlowVis*: a tool to visualize core signaling pathways associated with synergistic drug combinations

To visualize the results and gain a better understanding of the underlying mechanism of drug effects, we generated a core network of signaling interactions by applying a threshold. To enhance the interactions, we utilized the RShiny package and developed a visualization tool called *NetFlowVis* (The website link has been provided in the code availability part). This tool allows users to control the edge threshold and set a minimum number of nodes in each network component. Users have the option to select the desired cell line for visualization and mark specific signaling pathways for detailed analysis. As an example, **Figure 3b** demonstrates the

visualization of the cell line DU-145. Furthermore, by referring to the top 20 prior gene targets and their corresponding gene degree (node importance scores) for cell line DU-145, we observed that the majority of gene targets were included in the filtered signaling network interactions.

By utilizing the **NetFlowVis** visualization tool, users could take control of the size of the core signaling network. In the case of the DU-145 cell line, an edge threshold of 0.31 was selected, which corresponded to approximately the 98.5th percentile for edge weight among all edges. Additionally, a threshold of 23 was chosen for marking important nodes/genes in red, representing ~99.9th percentile for node importance scores across all nodes. To gain a deeper understanding of the underlying mechanism, the interactions between drugs and genes were analyzed, with the highest and lowest drug-gene interactions being represented by the colors purple and green, respectively. The analysis revealed that the drug **Paclitaxel** targeted the important gene **BCL2**, while the drug **Celecoxib** targeted less essential genes (see **Figure 3b**). This finding demonstrated that targeting the important genes selected by the **M3NetFlow** model led to higher drug scores, as depicted in **Figure 3c-f**. Moreover, by choosing a specific signaling pathway within the core signaling network, it was observed that the **IL-17 signaling pathway** was downstream of the critical gene **TRAF6**. Interestingly, experimental validation has shown the involvement of the **IL-17 signaling pathway** in promoting migration and invasion of **DU-145** cell lines through the upregulation of **MTA1** expression in prostate cancer⁶⁴.

Section E. Code Availability

M3NetFlow code: <https://github.com/FuhaiLiAiLab/M3NetFlow>

NetFlowVis app: <https://m3netflow.shinyapps.io/NetFlowVis/>

Acknowledgment

This study was partially supported by NIA R56AG065352 (to Li), NIA 1R21AG078799-01A1 (to Li/Province), NINDS 1RM1NS132962-01 (to Dickson/Marco/Cooper/Li), Children's Discovery Institute (CDI) M-II-2019-802, and NLM 1R01LM013902-01A1 to Li.

Author contributions

FL conceived this study. FL, HZ designed the model and prepared the manuscript. HZ implemented the model and analyzed the data and results. FL, HZ, PG, LD, WK, DD, RF, YC, PP discussed the method and results.

Declaration of interests

The authors declare no competing interests.

References

1. TCGA. <https://www.cancer.gov/about-nci/organization/ccg/research/structural-genomics/tcga>.
2. Barretina J, Caponigro G, Stransky N, et al. The Cancer Cell Line Encyclopedia enables predictive modelling of anticancer drug sensitivity. *Nature*. 2012;483(7391):603-607. doi:10.1038/nature11003
3. Yang W, Soares J, Greninger P, et al. Genomics of Drug Sensitivity in Cancer (GDSC): A resource for therapeutic biomarker discovery in cancer cells. *Nucleic Acids Res*. 2013;41(D1):955-961. doi:10.1093/nar/gks1111
4. Li F, Eteleeb A, Buchser W, et al. Weakly activated core inflammation pathways were identified as a central signaling mechanism contributing to the chronic neurodegeneration in Alzheimer's disease. *bioRxiv*. Published online 2021.
5. Li F, Oh I, Kumar S, et al. Loss of estrogen unleashing neuro-inflammation increases the risk of Alzheimer's disease in women. *bioRxiv*. Published online 2022:2022-2029.
6. TCGA. <https://www.cancer.gov/about-nci/organization/ccg/research/structural-genomics/tcga>.
7. Ghandi M, Huang FW, Jané-Valbuena J, et al. Next-generation characterization of the Cancer Cell Line Encyclopedia. *Nature*. 2019;569(7757):503-508. doi:10.1038/s41586-019-1186-3
8. Neff RA, Wang M, Vatansever S, et al. *Molecular Subtyping of Alzheimer's Disease Using RNA Sequencing Data Reveals Novel Mechanisms and Targets*. Vol 7.; 2021. <https://www.science.org>
9. Raghavachari N, Wilmot B, Dutta C. Optimizing Translational Research for Exceptional Health and Life Span: A Systematic Narrative of Studies to Identify Translatable Therapeutic Target(s) for Exceptional Health Span in Humans. *Journals of Gerontology -*

- Series A Biological Sciences and Medical Sciences*. 2022;77(11):2272-2280.
doi:10.1093/gerona/glac065
10. Hao N, Li Y, Jiang Y, et al. *Evolutionary-Conserved, Molecular Mechanisms Can Simultaneously Influence All Causes of Death. WADDINGTON'S LANDSCAPE OF CELL AGING DISRUPTION OF CPG ISLAND-MEDIATED CHROMATIN ARCHITECTURE AND TRANSCRIPTIONAL HOMEOSTASIS DURING AGING SESSION 1120 (SYMPOSIUM) THE LONGEVITY CONSORTIUM: MULTI-OMICS INTEGRATIVE APPROACH TO DISCOVERING HEALTHY AGING AND LONGEVITY DETERMINANTS Chair*.
11. Deelen J, Evans DS, Arking DE, et al. A meta-analysis of genome-wide association studies identifies multiple longevity genes. *Nat Commun*. 2019;10(1).
doi:10.1038/s41467-019-11558-2
12. Subramanian I, Verma S, Kumar S, Jere A, Anamika K. Multi-omics data integration, interpretation, and its application. *Bioinform Biol Insights*. 2020;14:1177932219899051.
13. Vaske CJ, Benz SC, Sanborn JZ, et al. Inference of patient-specific pathway activities from multi-dimensional cancer genomics data using PARADIGM. *Bioinformatics*. 2010;26(12):i237-i245.
14. Mihajlovićmihajlović K, Ceddia G, Malod-Dognin N, et al. Multi-omics integration of scRNA-seq time series data predicts new intervention points for Parkinson's disease. Published online 2023. doi:10.1101/2023.12.12.570554
15. Bodein A, Scott-Boyer MP, Perin O, Lê Cao KA, Droit A. TimeOmics: an R package for longitudinal multi-omics data integration. *Bioinformatics*. 2022;38(2):577-579.
doi:10.1093/bioinformatics/btab664
16. Kipf TN, Welling M. Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:160902907*. Published online 2016.
17. Hamilton WL, Ying R, Leskovec J. Inductive representation learning on large graphs. *Adv Neural Inf Process Syst*. 2017;2017-Decem(Nips):1025-1035.
18. Veličković P, Casanova A, Liò P, Cucurull G, Romero A, Bengio Y. Graph attention networks. *6th International Conference on Learning Representations, ICLR 2018 - Conference Track Proceedings*. Published online 2018:1-12.
19. Xu K, Hu W, Leskovec J, Jegelka S. How powerful are graph neural networks? *arXiv preprint arXiv:181000826*. Published online 2018.
20. Wang T, Shao W, Huang Z, et al. MOGONET integrates multi-omics data using graph convolutional networks allowing patient classification and biomarker identification. *Nat Commun*. 2021;12(1):3445.
21. Li X, Ma J, Leng L, et al. MoGCN: a multi-omics integration method based on graph convolutional network for cancer subtype analysis. *Front Genet*. 2022;13:806842.
22. Gao H, Zhang B, Liu L, Li S, Gao X, Yu B. A universal framework for single-cell multi-omics data integration with graph convolutional networks. *Brief Bioinform*. 2023;24(3):bbad081.
23. Rajadhyaksha N, Chitkara A. Graph Contrastive Learning for Multi-omics Data. *arXiv preprint arXiv:230102242*. Published online 2023.
24. Li G, Muller M, Thabet A, Ghanem B. DeepGCNs: Can GCNs go as deep as CNNs? *Proceedings of the IEEE International Conference on Computer Vision*. 2019;2019-Octob:9266-9275. doi:10.1109/ICCV.2019.00936
25. Cai C, Wang Y. A Note on Over-Smoothing for Graph Neural Networks. Published online 2020.
26. Morris C, Ritzert M, Fey M, et al. *Weisfeiler and Leman Go Neural: Higher-Order Graph Neural Networks*. www.aaai.org

27. Weisfeiler YuB, Leman A. The reduction of a graph to canonical form and the algebra which appears therein. *Nauchno-Technicheskaya Informatsia (NTI Series)*. 1968;2(9):12-16.
28. Chien E, Peng J, Li P, Milenkovic O. Adaptive universal generalized pagerank graph neural network. *arXiv preprint arXiv:200607988*. Published online 2020.
29. Nikolentzos G, Dasoulas G, Vazirgiannis M. k-hop graph neural networks. *Neural Networks*. 2020;130:195-205.
30. Abu-El-Haija S, Perozzi B, Kapoor A, et al. Mixhop: Higher-order graph convolutional architectures via sparsified neighborhood mixing. In: *International Conference on Machine Learning*. PMLR; 2019:21-29.
31. Brossard R, Frigo O, Dehaene D. Graph convolutions that can finally model local structure. *arXiv preprint arXiv:201115069*. Published online 2020.
32. Wang G, Ying R, Huang J, Leskovec J. Multi-hop attention graph neural network. *arXiv preprint arXiv:200914332*. Published online 2020.
33. Feng J, Chen Y, Li F, Sarkar A, Zhang M. *How Powerful Are K-Hop Message Passing Graph Neural Networks*.
34. Kanehisa MGS, Goto S. KEGG: kyoto Encyclopedia of Genes and Genomes. *Nucleic Acids Res*. 2000;28:27-30. doi:10.1093/nar/28.1.27
35. Oughtred R, Stark C, Breitkreutz BJ, et al. The BioGRID interaction database: 2019 update. *Nucleic Acids Res*. 2019;47(D1):D529-D541.
36. Stark C, Breitkreutz BJ, Reguly T, Boucher L, Breitkreutz A, Tyers M. BioGRID: a general repository for interaction datasets. *Nucleic Acids Res*. 2006;34(suppl_1):D535-D539.
37. Salton G. Modern information retrieval. (*No Title*). Published online 1983.
38. Preuer K, Lewis RPI, Hochreiter S, Bender A, Bulusu KC, Klambauer G. DeepSynergy: Predicting anti-cancer drug synergy with Deep Learning. *Bioinformatics*. 2018;34(9):1538-1546. doi:10.1093/bioinformatics/btx806
39. Zhang T, Zhang L, Payne P, Li F. Synergistic Drug Combination Prediction by Integrating Multi-omics Data in Deep Learning Models. *arXiv preprint arXiv:181107054*. Published online 2018.
40. Kipf TN, Welling M. Semi-supervised classification with graph convolutional networks. *5th International Conference on Learning Representations, ICLR 2017 - Conference Track Proceedings*. Published online 2017:1-14.
41. Veličković P, Casanova A, Liò P, Cucurull G, Romero A, Bengio Y. Graph attention networks. *6th International Conference on Learning Representations, ICLR 2018 - Conference Track Proceedings*. Published online 2018:1-12.
42. Shi Y, Huang Z, Feng S, Zhong H, Wang W, Sun Y. *Masked Label Prediction: Unified Message Passing Model for Semi-Supervised Classification*.
43. Corso G, Cavalleri L, Beaini D, Liò P, Veličković P, Deepmind V. *Principal Neighbourhood Aggregation for Graph Nets*.
44. Thekumparampil KK, Wang C, Oh S, Li LJ. *Attention-Based Graph Neural Network for Semi-Supervised Learning*.; 2018.
45. Sweatt JD. Mitogen-activated protein kinases in synaptic plasticity and memory. *Curr Opin Neurobiol*. 2004;14(3):311-317. doi:10.1016/j.conb.2004.04.001
46. Glass CK, Saijo K, Winner B, Marchetto MC, Gage FH. Mechanisms Underlying Inflammation in Neurodegeneration. *Cell*. 2010;140(6):918-934. doi:10.1016/j.cell.2010.02.016
47. Heneka MT, Carson MJ, El Khoury J, et al. Neuroinflammation in Alzheimer's disease. *Lancet Neurol*. 2015;14(4):388-405.
48. Munoz L, Ammit AJ. Targeting p38 MAPK pathway for the treatment of Alzheimer's disease. *Neuropharmacology*. 2010;58(3):561-568. doi:10.1016/j.neuropharm.2009.11.010

49. Malumbres M, Barbacid M. RAS oncogenes: the first 30 years. *Nat Rev Cancer*. 2003;3(6):459-465.
50. Bloom GS. Amyloid- β and tau: The trigger and bullet in Alzheimer disease pathogenesis. *JAMA Neurol*. 2014;71(4):505-508. doi:10.1001/jamaneurol.2013.5847
51. Braak H, Del Tredici K. The preclinical phase of the pathological process underlying sporadic Alzheimer's disease. *Brain*. 2015;138(10):2814-2833. doi:10.1093/brain/awv236
52. Talbot K. Brain insulin resistance in Alzheimer's disease and its potential treatment with GLP-1 analogs. *Neurodegener Dis Manag*. 2014;4(1):31-40. doi:10.2217/nmt.13.73
53. Zhao Y, Yang J, Liao W, et al. Cytosolic FoxO1 is essential for the induction of autophagy and tumour suppressor activity. *Nat Cell Biol*. 2010;12(7):665-675.
54. Poo M ming. Neurotrophins as synaptic modulators. *Nat Rev Neurosci*. 2001;2(1):24-32.
55. Bomfim TR, Forny-Germano L, Sathler LB, et al. An anti-diabetes agent protects the mouse brain from defective insulin signaling caused by Alzheimer's disease-associated A β oligomers. *J Clin Invest*. 2012;122(4):1339-1353.
56. Nistri S, Bani D. Relaxin in vascular physiology and pathophysiology: possible implications in ischemic brain disease. *Curr Neurovasc Res*. 2005;2(3):225-233.
57. O'Neil J, Benita Y, Feldman I, et al. An unbiased oncology compound screen to identify novel combination strategies. *Mol Cancer Ther*. 2016;15(6):1155-1162. doi:10.1158/1535-7163.MCT-15-0843
58. Zagidullin B, Aldahdooh J, Zheng S, et al. DrugComb: An integrative cancer drug combination data portal. *Nucleic Acids Res*. 2019;47(W1):W43-W51. doi:10.1093/nar/gkz337
59. Van Der Meer D, Barthorpe S, Yang W, et al. Cell Model Passports - a hub for clinical, genetic and functional datasets of preclinical cancer models. *Nucleic Acids Res*. 2019;47(D1):D923-D929. doi:10.1093/nar/gky872
60. Barretina J, Caponigro G, Stransky N, et al. The Cancer Cell Line Encyclopedia enables predictive modelling of anticancer drug sensitivity. *Nature*. 2012;483(7391):603-607. doi:10.1038/nature11003
61. Ogata H, Goto S, Sato K, Fujibuchi W, Bono H, Kanehisa M. KEGG: Kyoto encyclopedia of genes and genomes. *Nucleic Acids Res*. Published online 1999:28. doi:10.1093/nar/27.1.29
62. Wishart DS, Feunang YD, Guo AC, et al. DrugBank 5.0: A major update to the DrugBank database for 2018. *Nucleic Acids Res*. 2018;46(D1):D1074-D1082. doi:10.1093/nar/gkx1037
63. Holbeck SL, Camalier R, Crowell JA, et al. The National Cancer Institute ALMANAC: A comprehensive screening resource for the detection of anticancer drug pairs with enhanced therapeutic activity. *Cancer Res*. 2017;77(13):3564-3576. doi:10.1158/0008-5472.CAN-17-0489
64. Guo N, Shen G, Zhang Y, Moustafa AA, Ge D, You Z. Interleukin-17 promotes migration and invasion of human cancer cells through upregulation of MTA1 expression. *Front Oncol*. 2019;9(JUN):1-11. doi:10.3389/fonc.2019.00546