

Systems biology

mosGraphGen: a novel tool to generate multi-omics signaling graphs to facilitate integrative and interpretable graph AI model development

Heming Zhang ¹, Dekang Cao¹, Zirui Chen¹, Xiuyuan Zhang¹, Yixin Chen², Cole Sessions³, Carlos Cruchaga^{4,5,6}, Philip Payne ¹, Guangfu Li⁷, Michael Province⁸, Fuhai Li ^{1,3,6,*}

¹Institute for Informatics, Data Science and Biostatistics (I2DB), Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

²Department of Computer Science and Engineering, Washington University in St. Louis, Saint Louis, MO 63130, United States

³Department of Pediatrics, Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

⁴Department of Psychiatry, Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

⁵Hope Center for Neurological Disorders, Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

⁶NeuroGenomics and Informatics Center, Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

⁷Department of Surgery, School of Medicine, University of Connecticut, Farmington, CT 06030, United States

⁸Division of Statistical Genomics, Department of Genetics, Washington University School of Medicine, Saint Louis, MO 63110-1010, United States

*Corresponding author. Institute for Informatics, Data Science and Biostatistics (I2DB), Washington University School of Medicine, 660 S. Euclid Ave., Saint Louis, 63110-1010, MO, USA; Department of Pediatrics, Washington University School of Medicine, 660 S. Euclid Ave., Saint Louis, 63110-1010, MO, USA; NeuroGenomics and Informatics Center, Washington University School of Medicine, 660 S. Euclid Ave., Saint Louis, 63110-1010, MO, USA.
E-mail: fuhai.li@wustl.edu

Associate Editor: Shanfeng Zhu

Abstract

Motivation: Multi-omics data, i.e. genomics, epigenomics, transcriptomics, proteomics, characterize cellular complex signaling systems from multi-level and multi-view and provide a holistic view of complex cellular signaling pathways. However, it remains challenging to integrate and interpret multi-omics data for mining critical biomarkers. Graph AI models have been widely used to analyze graph-structure datasets, and are ideal for integrative multi-omics data analysis because they can naturally integrate and represent multi-omics data as a biologically meaningful multi-level signaling graph and interpret multi-omics data via graph node and edge ranking analysis. Nevertheless, it is nontrivial for graph-AI model developers to pre-analyze multi-omics data and convert the data into biologically meaningful graphs, which can be directly fed into graph-AI models.

Results: To resolve this challenge, we developed mosGraphGen (multi-omics signaling graph generator), generating Multi-omics Signaling graphs (mos-graph) of individual samples by mapping multi-omics data onto a biologically meaningful multi-level background signaling network with data normalization by aggregating measurements and aligning to the reference genome. With mosGraphGen, AI model developers can directly apply and evaluate their models using these mos-graphs. In the results, mosGraphGen was used and illustrated using two widely used multi-omics datasets of The Cancer Genome Atlas (TCGA) and Alzheimer's disease (AD) samples.

Availability and implementation: The code of mosGraphGen is open-source and publicly available via GitHub: <https://github.com/FuhaiLiLab/mosGraphGen>.

1 Introduction

Along with the advancements of next-generation sequencing (NGS) and high-throughput technologies, multi-omics datasets including genomics, epigenomics, transcriptomics, proteomics, and metabolomics, have been abundantly generated. The multi-omics datasets characterize the dysfunctional molecular mechanisms of complex diseases from different aspects. The integrative analysis of multi-omics data can provide a holistic view of the mysterious and complex signaling systems of cells from complex diseases. Multi-omics data-driven studies are at the forefront of precision medicine and healthcare, which are critical for uncovering novel

therapeutic targets and discovering novel and effective drugs and cocktail treatments. Recently, a new Multi-omics for Health and Disease Consortium was established by the National Institutes of Health (NIH), aiming to advance multi-omics data generation and integrative analysis for human disease and health research.

Large-scale multi-omics datasets of different diseases have been generated and publicly accessible. For example, multi-omics data of >40 000 cancer samples from ~69 primary sites were generated in The Cancer Genome Atlas (TCGA) (Grossman *et al.* 2016, Sanchez-Vega *et al.* 2018) project, which is a comprehensive initiative launched by the National

Cancer Institute (NCI) and the National Human Genome Research Institute (NHGRI). The datasets were generated and publicly accessible (Grossman *et al.* 2016, Goldman *et al.* 2020) to characterize and understand the genomic alterations and molecular profiles of >30 cancer types/subtypes. In addition, to increase the range of characterizing the cancer cell lines, the Cancer Cell Line Encyclopedia (CCLE) (Barretina *et al.* 2012) study provides multi-omics datasets of ~1000 cell lines for 36 tumor types. The datasets have been widely used to study genetic biomarkers and associations with drug and cocktail responses. Also, the multi-omics datasets of Alzheimer's disease (AD) multi-center cohorts, e.g. Mayo Clinic (Bennett *et al.* 2018, De Jager *et al.* 2018), Mount Sinai/JJ Peters VA Medical Center Brain Bank (MSBB-AD) (Allen *et al.* 2016), and ROSMAP (Neff *et al.* 2021), have been generated and publicly accessible via the synapse website of AD-AMP (Greenwood *et al.* 2020, Li *et al.* 2022a). Moreover, large-scale multi-omics data, like genetics, epigenetics, transcriptomics, proteomics, and metabolomics (Deelen *et al.* 2019, Raghavachari *et al.* 2022) are being generated by ongoing NIA-supported exceptional longevity (EL) studies, including the Long-Life Family Study (LLFS), the Longevity Consortium (LC), the longevity genomics (LG), and the Integrative Longevity Omics (ILO), to understand and identify protective factors/targets, biological processes, and signaling pathways that promote health and life span in exceptionally long-lived individuals. These datasets are valuable for studying the complex molecular mechanisms of complex diseases. Moreover, single-level multi-omics datasets are generated. The advance in single-cell (sc) omics has made it a powerful technology to investigate the genetic and functional heterogeneity of diverse cell types within disease microenvironment/niche (Lee *et al.* 2020, Baysoy *et al.* 2023). Compared with the tissue-level, sc omics datasets provide a finer view of the complex signaling system within diverse cell types such as disease and immune cell types, subtypes, and different cell states (Lee *et al.* 2020, Baysoy *et al.* 2023).

Multi-omics data integration and interpretation remain an open problem, and network-based models are ideal for multi-omics data analysis. Analyzing and interpreting multi-omics data is complex and computationally challenging. It requires novel and sophisticated bioinformatics and data integration techniques to uncover meaningful insights. The following section summarizes related works of multi-omics data analysis and provides a comprehensive review of existing multi-omics data integration analysis models (Subramanian *et al.* 2020). Specifically, these models were divided into a few categories: similarity, correlation, Bayesian, multivariate, fusion, and network-based models (Subramanian *et al.* 2020). Network-based models, such as PARADIGM (Sedgewick *et al.* 2013) (Pathway Representation and Analysis by Direct Inference on Graphical Models), are some of the most widely used methods; Visible Neural Network [VNN (Yu *et al.* 2018)] models [e.g. DCell (Ma *et al.* 2018) and DrugCell (Kuenzi *et al.* 2020)] were proposed using large hierarchical deep learning architecture to model the hierarchical organization of biological processes and to predict drug response with important biomarkers. However, the signaling cascade level activity has not been specifically investigated in DCell (Ma *et al.* 2018) and DrugCell (Kuenzi *et al.* 2020) models. The recently developed graph neural network (GNN) (van den Berg *et al.* 2018) model is ideal for graph-structure data

analysis tasks and multi-omics data analysis with its ability to (i) integrate multi-omics data as features of nodes and (ii) model the signaling interactions or signaling cascades among the proteins using protein-protein interactions (Szklarczyk *et al.* 2015, Oughtred *et al.* 2019) or signaling pathways (Kanehisa and Goto 2000, Slinger *et al.* 2018); (iii) the key signaling target identification and pathway analysis can be achieved via node and edge ranking in GNN (Dong *et al.* 2023, Zhang *et al.* 2024). Here are a few recent studies that made use of the network-/graph-based analysis for multi-omics data analysis. MOGONET (Wang *et al.* 2021) used the omics-specific similarity graphs among samples and used the GCN (Lee *et al.* 2020) to learn patients' labels from each omics data independently. The MoGCN (Li *et al.* 2022b) employed a similar design using the patient similarity. GCN-SC (Gao *et al.* 2023) employed GCN to combine single-cell multi-omics data. MOGCL (Rajadhyaksha and Chitkara 2023) explored graph contrastive learning to pretrain the GCN on the multi-omics dataset. However, these models have yet to incorporate signaling pathways.

Computational tools that can generate graph-structure data for individual samples, i.e. mapping multi-omics data onto a biologically meaningful background signaling graph, are urgently needed and critical for developing graph-AI models. As introduced above, we found the top two challenges in graph-AI for multi-omics data analysis are: building a large-scale biologically meaningful background signaling graph to integrate multi-omics data and developing effective approaches to rank key signaling targets and pathways from the large-scale background signaling graph. For many graph-AI developers, particularly those without training in multi-omics data pre-analysis and integrative analysis, it is a challenging task to map the multi-omics data of individual samples onto a meaningful network and generate graphs as inputs ready for graph-AI models. This is because omics-specific pre-analyses are needed to convert the raw data into the standard annotation, like multi-level gene-protein-promoter-enhancer-associations, for multi-omics data integration; and it is essential to link these annotations via biologically meaningful interactomes, including protein-protein interactions, signaling pathways, and transcription factor (TF)-target interactions. To address these challenges, we developed mosGraphGen, a multi-omics signaling graph generator that converts multi-omics datasets into graph-structured data, which can then be used as inputs for graph-AI models. This facilitates integrative and interpretable multi-omics data analysis, such as target and edge ranking analysis within graph-AI models.

Open-source: The code of mosGraphGen is open-source and publicly available via GitHub: <https://github.com/FuhaiLiAiLab/mosGraphGen>. The following sections provide the details of the methodology and results.

2 Methodology and materials

2.1 Multi-omics data pre-processing

Figure 1 presents a detailed schematic of the pipeline designed to transform extensive multi-omics datasets into graph-based models. This pipeline encapsulates the sequential and methodical approach employed to preprocess, integrate, and ultimately represent data from multi-omics datasets, clinical dataset, reference genome dataset, and regulatory network dataset as coherent networks. This pipeline ensures the systematic standardization with aggregating and resolving duplicates across multi-level and integration necessary for

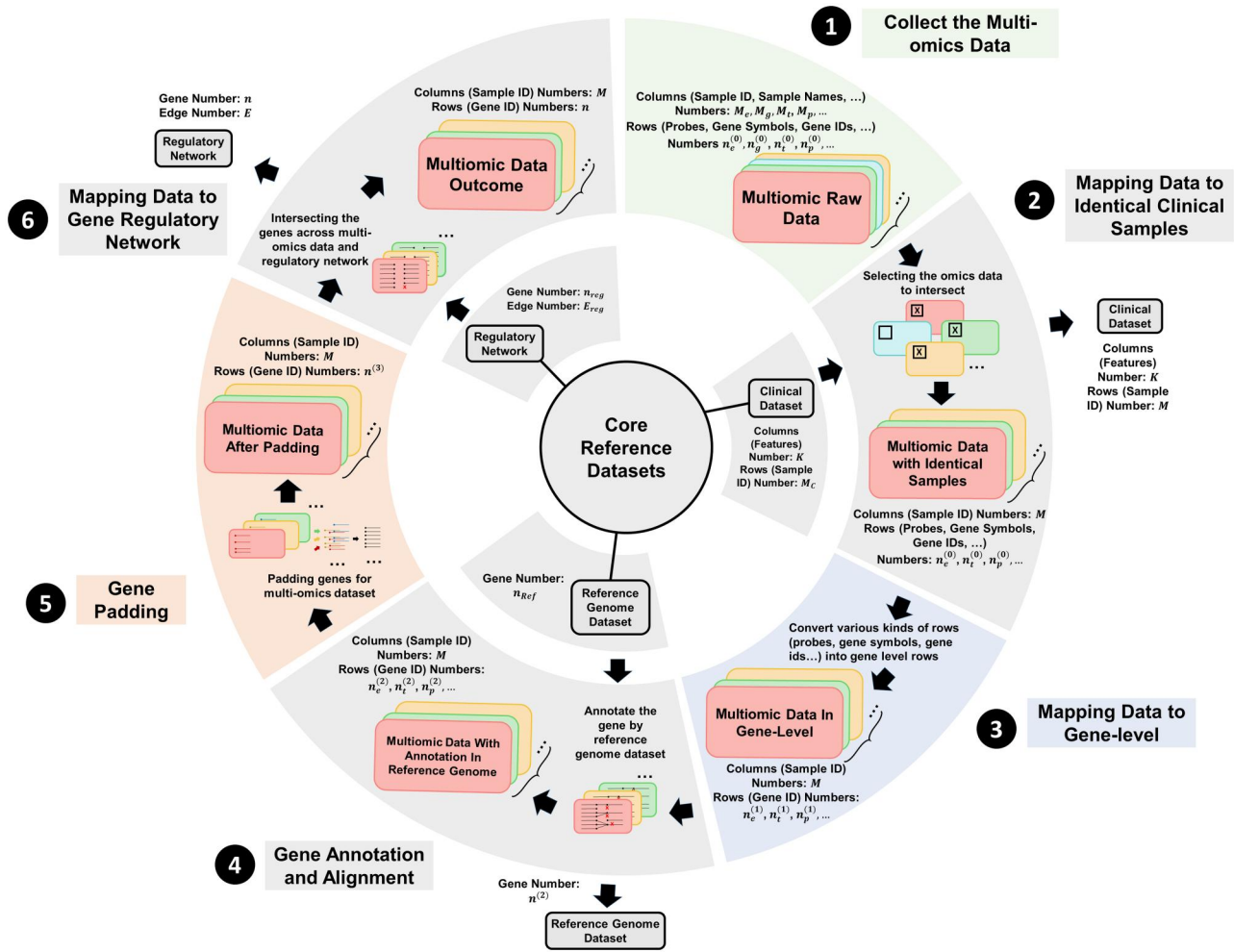


Figure 1. Flowchart multi-omics data processing. Collect various kind of multi-omics data like (epigenomics, genomics, transcriptomics, proteomics, etc.). Integrating the multi-omics data with clinical data and getting the identical samples across the datasets. Converting the rows (probes, gene symbols, gene ids, etc.) into gene-level by aggregating the same measurements for one gene or by dropping the duplicates for gene synonym. Aligning genes by reference genome so that the final annotation for each gene in multi-omics data will be generated. Unifying the number of genes across multi-omics datasets and make the data imputation by filling zero values in empty spaces. Integrating the gene regulatory network with multi-omics datasets and generating the final multi-omics data with identical number of samples and genes.

generating multi-omics datasets that are critical for advancing our understanding of complex biological systems by referring to the genome dataset. The specific procedures within the pipeline are outlined in the following description.

2.2 Multi-omics datasets

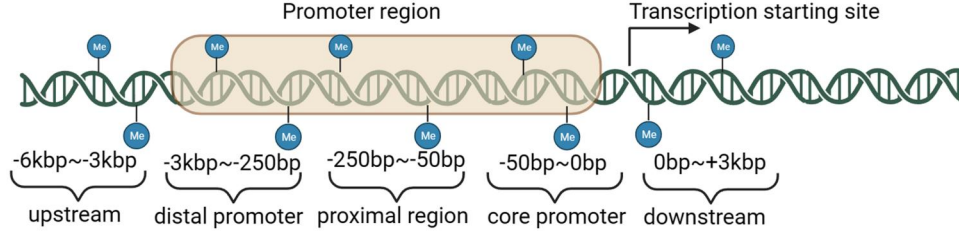
The multi-omics data can be downloaded from multiple public available datasets. For example, multi-omics data and their related datasets for the TCGA dataset and ROSMAP dataset can be downloaded from public available dataset. After downloading the multi-omics (epigenomics, genomics, transcriptomics, proteomics, etc.) datasets from resources, the multi-omics datasets will be converted into 2D spaces data frames with columns (sample IDs, sample names, etc.) of $M_e, M_g, M_t, M_p, \dots$ and rows (probes, gene symbols, gene IDs, etc.) of $n_e^0, n_g^0, n_t^0, n_p^0, \dots$. For example, the methylation data in the UCSC dataset have row type of CpG probes target IDs and n_e^0 is the number of rows in this methylation dataset. And column type will be sample IDs and M_e is the number of columns in the UCSC methylation dataset.

2.3 Align multi-omics data with clinical information

By selecting the multi-omics datasets from the data downloaded, the clinical dataset with M_c samples and K features will be integrated into the multi-omics data. In these clinical datasets, information of each sample [like OS (overall survival), PFI (progression-free interval), DSS (disease-specific survival), DFI (disease-free interval), vital status in UCSC dataset or sex, age_at_visit, ceradsc, cogdx in ROSMAP], will be used as the labels of each sample for different tasks. Since different omics data might use different sample information (e.g. OS binary classification value will be used as the label to indicate the survival status for samples in the UCSC dataset and ceradsc types will be used as the label to indicate the AD level for each patient in ROSMAP study), the step will associate each omics data to the unique sample ID, which associate with clinical information. Nevertheless, the multi-omics datasets in ROSMAP have a sample ID recognition problem. Hence, the sample ID mapping (e.g. PT-M5AF to R1822146 in ROSMAP dataset) system across multi-omics datasets should be constructed to align identical samples across datasets (check Table 1 for details).

Table 1. Mapping relations example for samples in ROSMAP.

individualID	mirna_id	mwas_id	mrna_id	cnvdata_id	projid	Study
R1822146	b01.127N	PT-M5AF	575_120521_2	SM-CTDQQ	20271359	ROS
R3143439	b01.128N	PT-BZBT	502_120515_1	SM-CTEMJ	38967303	MAP
R6879714	b01.130C	PT-3PTN	660_120530_1	SM-CTEEU	65736039	MAP
R8963331	b01.130N	PT-M5GT	607_120523_2	SM-CJEKL	20197364	ROS
...	...	..	...	...	...	...

**Figure 2.** Dividing methylation regions into five regions.

However, only a part of sample IDs in different omics data are overlapped, which will reduce the M_C samples to an identical number of samples with M in the selected datasets (e.g. epigenomics, transcriptomics proteomics, etc.), depending on the omics data selection strategies. Also, after the intersection with the clinical datasets, the number of rows in selected datasets (e.g. epigenomics, transcriptomics proteomics, etc.) will be n_e^0 , n_t^0 , n_p^0 , etc.

2.4 Mapping data to gene-level

Given that the rows of each multi-omics dataset vary, the methods used for mapping the rows to gene level are unique for each dataset. In detail, each omics data has their special mapping steps described as follows.

2.4.1 Epigenomics/methylation data pre-analysis

The methylation data will be retrieved from the UCSC XENA DNA methylation (450k) for the UCSC dataset and ROSMAP_arrayMethylation_imputed for ROSMAP dataset. However, the rows in the above methylation datasets are probes to detect the methylation condition in each individual CpG site. To convert the CpG site level methylation to TSS level, the GEO GPL16304 Platform will be considered for both datasets, which provides annotations, such as Distance_closest_TSS and Closest_TSS_gene_name, for each probe in the methylation (450k) dataset, mapping probe IDs to their corresponding gene names (e.g. cg00001583 to NR5A2). Furthermore, the gene name from the GPL16304 file was replaced with the probe ID in the methylation data by merging with the probe ID file. Hence, the nearest TSS gene name was incorporated into the new data frame.

In the next step, a preliminary analysis is essential to convert the TSS level methylation to gene level, focusing specifically on the -6 kb to $+3$ kb range around transcription start sites (TSS). This range will be subdivided into five distinct regions and probes located outside this specified range will be excluded from the analysis, which will convert rows from the CpG site level to the gene-level (Butler and Kadonaga 2002, Saxonov *et al.* 2005, Saintenac *et al.* 2011, Duttke *et al.* 2015, Vanaja and Yella 2022). In detail, promoter regions were delineated into three categories: The core promoter region (0 – 50 bp upstream of TSS), the proximal promoter

region (50 – 250 bp upstream of TSS), and the distal promoter region (250 – 3000 bp upstream of TSS). Additionally, the area 3000 – 6000 bp upstream of the distal promoter region was defined as the upstream region, while the downstream region encompassed the area from the TSS extending to 3000 bp downstream (check Fig. 2 for division of promoter and whole methylation data regions).

According to the defined regions of interest for methylation values, the TSS level methylation will be transformed into the gene level by using mean aggregation and rows with TSS closest distances outside the above regions will be omitted. This resulted in five CSV files, each containing the average methylation levels for genes in distinct five regions as mentioned before. To keep the dimensions of methylation data in each region in the same format, the unique genes in all the regions should be integrated and padding the empty value with zero and the number of rows in methylation data will be $n_e^{(1)}$ for UCSC or ROSMAP dataset.

2.4.2 Genetic variations pre-analysis

For UCSC genome variant datasets, the preprocessed genomic dataset is provided in gene-level mutation with binary value (0 wild-type for and 1 for nonsilent mutation). To remove the duplicated rows named with same gene, mean values was applied to aggregate the mutations. With respect to Copy Number Variations data in ROSMAP dataset, we employed the pyensembl package to assign gene names based on the chromosomal start and end positions (e.g. Start-End 830676–834492 to gene LINC01128). Given multiple types of variation are defined as deletions (DEL), duplications (DUP), and multiple copy number variations (mCNV), genome variants data were summed based on multiple genetic variation by variation types to capture the cumulative effect of variations of the same gene. In contrast to other multi-omics data, to keep the dimensions of genetic variation data in each type of variation in same format, the unique genes in all the variation types should be integrated and padding the empty value with -1 to maintain the integrity of the variation dataset. In the end, the output dataset after this processing steps will have the number of rows as $n_g^{(1)}$ for UCSC or ROSMAP dataset.

2.4.3 Transcriptomics data pre-analysis

For UCSC RNA sequencing datasets, gene names were directly converted using an ID mapping table associated with the gene expression dataset from the UCSC website. With regard to ROSMAP RNA sequencing data, we removed the version suffix from gene IDs in the RNA sequencing data (e.g. transforming ENSG00000167578.11 to ENSG00000167578) to adapt version updates in annotations. Using the official Ensembl dataset as the reference, our sequencing data was refined by replacing gene IDs with their respective gene names based on the established correspondence between gene IDs and names within the Ensembl resource. Mean values were calculated for duplicate gene symbols to ensure a single, representative value for each gene for both the UCSC and ROSMAP transcriptomics dataset, resulting in the number of rows as $n_t^{(1)}$.

2.4.4 Proteomic data pre-analysis

Regarding proteomics data in the ROSMAP dataset, we first discarded protein accession numbers, focusing solely on gene/protein names (e.g. VAMP1|P23763 to VAMP1) to ensure consistency with the gene-centric analysis framework of mosGraphGen. In contrast to other omics datasets, proteomics data in both UCSC and ROSMAP datasets did not require a specific pre-analysis and mapping phase due to its direct association with gene/protein names. Mean values were calculated for duplicated genes to ensure a single, representative value for each gene, ending with the number of rows as $n_p^{(1)}$.

2.4.5 Gene annotations, alignment, and padding

After getting the gene-level data for selected multi-omics datasets, the reference genome dataset (e.g. Ensembl dataset, containing n_{ref} genes) will be chosen as the standard for annotating the genes across multi-omics datasets. Therefore, a comprehensive comparison was conducted with the Ensembl dataset to ensure the coherence and accuracy of our genome variant dataset. Genes identified in pre-processed multi-omics datasets that were not present in the Ensembl dataset were excluded from further analysis to ensure that only genes with recognized genomic annotations were included. Therefore, the selected multi-omics datasets will contain M identical samples and rows with $n_e^{(2)}$, $n_t^{(2)}$, $n_p^{(2)}$, ... To unify the number of genes across multi-omics datasets, padding methods will be applied here with integrating unique genes in all datasets. Hence, zero values will be filled in the empty spaces to complete this data imputation method. Specifically, genomics data will be filled with -1 for empty spaces to distinguish the empty values with the nonmutation conditions. After this padding step, the number of genes will be identical across all multi-omics datasets with $n^{(3)}$ rows.

2.4.6 Mapping data to gene regulatory network

To integrate the gene regulatory network for multi-omics data, the selected gene regulatory network [e.g. KEGG (Ogata *et al.* 1999, Kanehisa and Goto 2000), BioGRID (Stark *et al.* 2006, Oughtred *et al.* 2019), or STRING (Szklarczyk *et al.* 2019)] with n_{reg} genes and E_{reg} protein-protein interactions will be intersected with the multi-omics datasets. Filtering out the genes not overlapped with the gene regulatory network, the multi-omics datasets will contain an identical number of samples M and genes n .

3 Results

3.1 Overview of mosGraphGen

From a computational perspective, integrating multi-omics data into the graph model involved several steps. First, multi-omics data from each patient were collected. In the meanwhile, the related datasets such as clinical datasets as sample labels, reference genome dataset as gene annotation, and gene regulatory network as organizing the gene knowledge graph will also downloaded and integrated into multi-omics dataset as mentioned in the method part of multi-omics data pre-processing. However, most of the networks ignored the details of biological processes such as methylation and transcription. Hence, the detailed version was depicted in Fig. 3 with organizing the promoter regions and transcription activity. Therefore, the re-organized network could be leveraged as the universal knowledge graph for each patient as the graph model. The input features for each patient were generated by integrating the matrices from each omics feature. Utilizing the network and feature inputs, graph neural network models were trained to predict patient outcomes and identify patient-specific critical biomarkers. This was achieved by extracting attention scores or edge weights from the trained models. The complete processes, including data loading, model training, and downstream analysis, are publicly accessible on GitHub documents: <https://github.com/FuhaiLiAiLab/mosGraphGen/blob/main/README.md>

3.2 Graphical overview of mosGraphGen

Figure 3 shows the architecture of the graph neural network model. The model has input $\mathcal{X} = \{X^{(1)}, X^{(2)}, \dots, X^{(m)}, \dots, X^{(M)}|A\}$, ($X^{(m)} \in \mathbb{R}^{n \times d}$, $A \in \mathbb{R}^{n \times n}$), where M denotes the number of samples/patients in each input dataset, and n denotes the number of nodes/genes used in the graph AI model; A is an adjacency matrix for the number of n nodes for the generated multi-omics signaling graph; and d is the number of multi-omics features organized in the initial embeddings in each node. To predict the outcome of all samples $\mathcal{Y}$ ($\mathcal{Y} \in \mathbb{R}^{M \times C}$), where C is the number of classes and $\mathcal{Y} = \{y^{(1)}, y^{(2)}, \dots, y^{(m)}, \dots, y^{(M)}\}$, we build a machine-learning model $f(\cdot)$ with $f(\mathcal{X}, A) = \mathcal{Y}$, where $\mathcal{X}$ denotes all of the data points in the dataset.

3.2.1 Graph-AI ready data format

For both UCSC and ROSMAP datasets, the processed datasets were saved in the special folder in CSV format firstly. Afterwards, the clinical datasets were considered as the graph prediction labels, i.e. the $\mathcal{Y}$ ($\mathcal{Y} \in \mathbb{R}^{M \times C}$) in the methodology part. Another task was to combine the multi-omics features in each CSV file by concatenating those features into a 3D NumPy array, i.e. $\mathcal{X}$ in the math part, which contained information about each multi-omics feature in all genes for every sample/patient. Since the torch-geometric package was used here to enhance the computing speed and save more spaces in GPU, the adjacency matrix was transformed into the edge index format with the dimension of $(E, 2)$, where E was the number of edges in the proposed graph model or the number of edges/links in the gene regulatory network. In addition, to validate the proposed model, a 5-fold cross-validation was used. In total, there were 3592 model input data points for the UCSC dataset and 138 model input data points for the ROSMAP dataset if all multi-omics datasets are selected. In each fold model, 4 folds of the data points were used as the

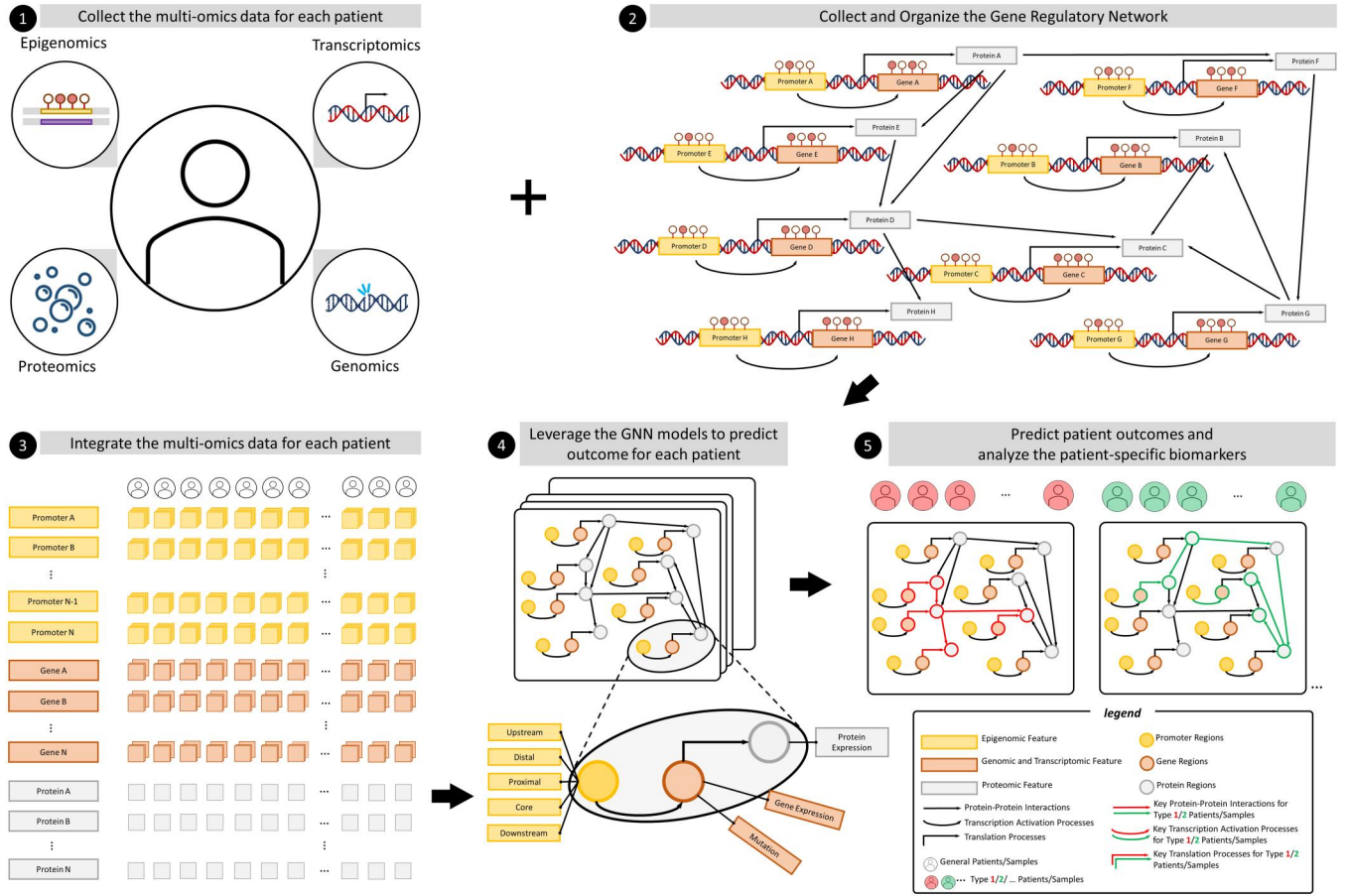


Figure 3. Overview of mosGraphGen. Multi-omics dataset (epigenomics, transcriptomics, genomics, and proteomics) for each patient will be collected. The gene regulatory network from KEGG, BioGRID, or STRING will be downloaded and integrate the protein–protein interactions with the promoters and transcriptions process. Hence, the re-organized regulatory network will incorporate more details and be formalized as the input matrix for graph model. The collected multi-omics dataset will be processed and integrate into the matrices format for each patient. The well-formatted matrices (features and regulatory network) will be used as the input for various graph neural network models to make the prediction for each patient. Downstream analysis will be made by extracting the edge weights from the attention trained from the selected graph neural network. Hence, the patient-specific analysis of discovering the critical biomarkers will be generated.

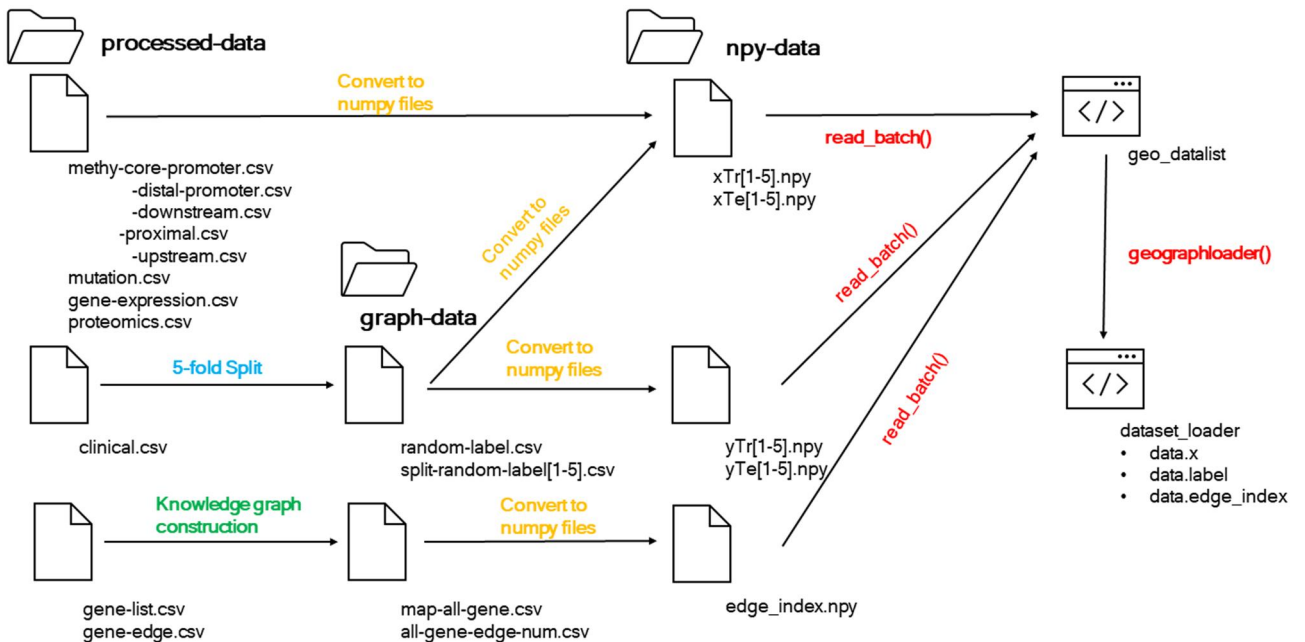


Figure 4. Loading processed data into graph-AI format.

Table 2. UCSC Database resources.

Database	Description	Link
UCSC Xena DNA methylation (450k)	DNA methylation dataset generated using the Illumina Infinium HumanMethylation450 BeadChip array.	https://xenabrowser.net/datapages/?dataset=jhu-usc.edu_PANCAN_HumanMethylation450.betaValue_whitelisted.tsv.synapse_download_5096262.xena&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
Protein expression—RPPA	Quantifying protein expression	https://xenabrowser.net/datapages/?dataset=TCGA-RPPA-pancan-clean.xena&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
Somatic mutation (SNP and INDEL)—Gene level nonsilent mutation	The TCGA Unified Ensemble “MC3” gene-level mutation dataset identifies somatic mutations in various cancers, marking nonsilent mutations (1) that alter protein sequences and wild type 0 for no mutations.	https://xenabrowser.net/datapages/?dataset=mc3.v0.2.8.PUBLIC.nonsilentGene.xena&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
Gene expression RNAseq—TOIL RSEM fpkm	gene expression data derived from RNAseq, processed using the TOIL pipeline, with expression levels estimated using RSEM and normalized as FPKM values.	https://xenabrowser.net/datapages/?dataset=tcga_RSEM_gene_fpkm&host=https%3A%2F%2Ftoil.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
GEO GPL16304 Platform	Illumina HumanMethylation450 BeadChip [UBC enhanced annotation v1.0]	https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GPL16304
Curated clinical data	Contains patient clinical features, achieved from paper “An Integrated TCGA Pan-Cancer Clinical Data Resource (TCGA-CDR) to drive high quality survival outcome analytics.”	https://xenabrowser.net/datapages/?dataset=Survival_SupplementalTable_S1_20171025_xena_sp&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
Immune subtype	Model-based immune subtype	https://xenabrowser.net/datapages/?dataset=Subtype_Immune_Model_Based.txt&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
Molecular subtype	Phenotype data	https://xenabrowser.net/datapages/?dataset=TCGASubtype_20170308.tsv&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443
sample type and primary disease	sample type and primary disease information combined from all individual TCGA cohorts	https://xenabrowser.net/datapages/?dataset=TCGA_phenotype_denseDataOnlyDownload.tsv&host=https%3A%2F%2Fpancanatlas.xenahubs.net&removeHub=https%3A%2F%2Fxena.treehouse.gi.ucsc.edu%3A443

Table 3. Dimensions of the UCSC multi-omics data before and after the data processing.

Database	Before process	After
Methylation dataset	(396 065 probes, 9664 samples)	
Mutation dataset	(40 543 genes, 9104 samples)	(2117 genes in KEGG,
RNASeq dataset	(60 498 gene IDs, 10 535 samples)	16 539 genes in BioGRID,
Protein expression dataset	(258 genes, 7754 samples)	15 923 genes in STRING,
		<i>M</i> samples)
Processed clinical dataset	(6385 samples, 44 features)	(22 features, <i>M</i> samples)

training dataset, and the rest 1-fold was used as the testing dataset. Finally, the processed data with NumPy format ($\mathcal{X}$, $\mathcal{Y}$, E) were loaded with DataLoader class in torch-geometric to transform them into the machine-readable CUDA tensor (check Fig. 4 for details).

3.2.2 Multi-omics datasets of TCGA cancer samples

Table 2 presents download links for multi-omics datasets of TCGA cancer samples, encompassing analyses of DNA methylation patterns, somatic mutations—including Single Nucleotide Polymorphisms (SNPs) and Insertions/Deletions (INDELs)—as well as gene expression profiles obtained through RNA sequencing. As to methylation data, the conversion from CpG site-level data to gene-level data will be

conducted based on the aforementioned method. The detailed data processing results are described as follows.

As detailed in the method above, raw multi-omics datasets in UCSC are collected and loaded into the pandas dataframe (check Table 3 for rows and columns information). The methylation dataset encompasses 396 065 probes, targeting the CpG sites of genes across. Since the regions of interest around TSS were predefined as mentioned in Fig. 2, rows with TSS distances that fell outside these areas were excluded and each CpG probe was assigned with corresponding TSS closest gene name, reducing the number of rows from 396 065 to 242 313. By splitting the methylation files into 5 files according to predefined promoter regions, the number of rows (genes) will be 5638, 17 797, 16 169, 10 139, and 21 757 for the upstream region, distal promoter region,

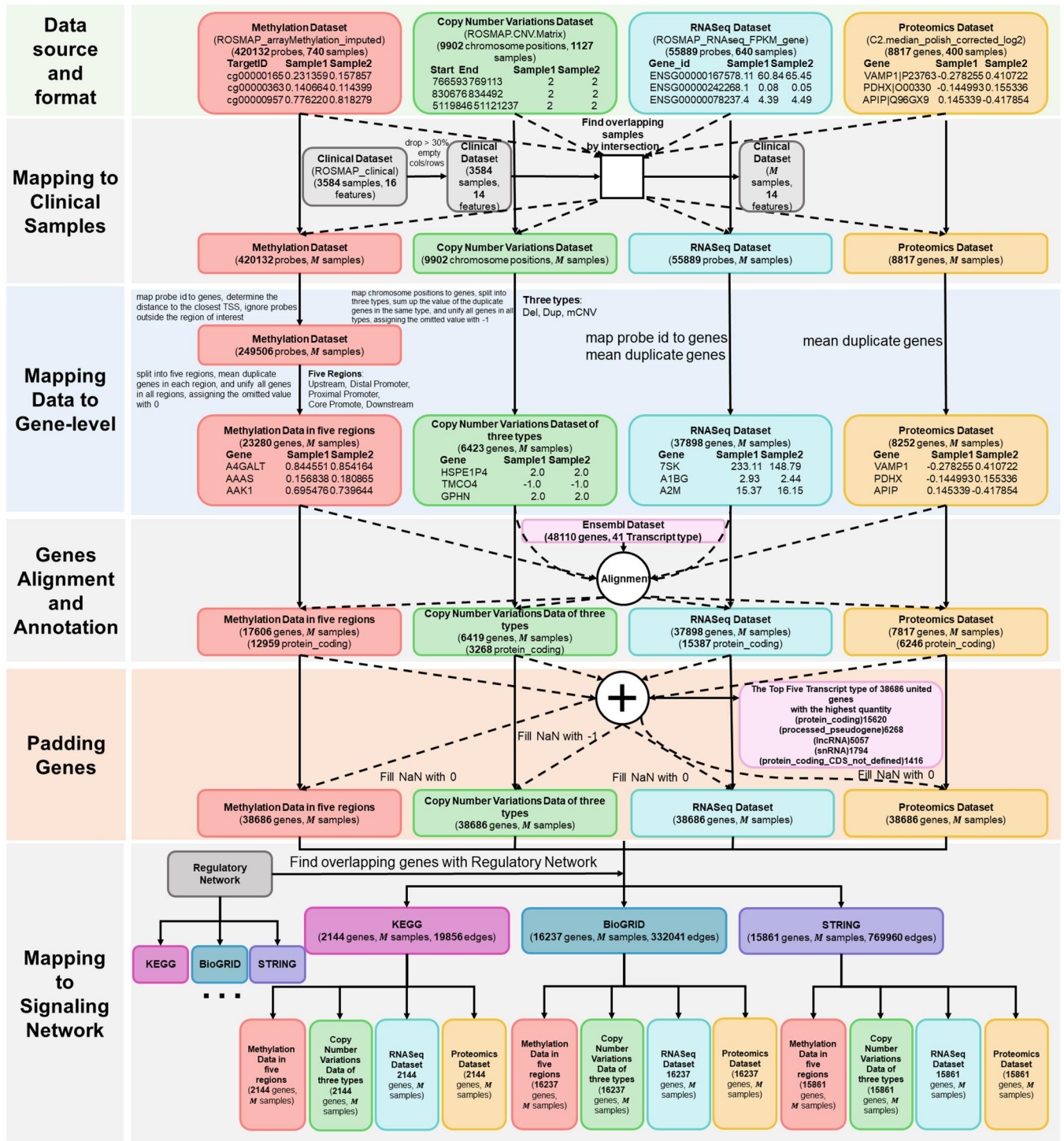


Figure 5. Flowchart of TCGA cancer multi-omics data analysis. The methylation dataset undergoes probe-to-gene mapping, with a focus on the -6 kb to +3 kb range around the TSS, followed by division into five genomic regions. The Protein Expression Dataset includes 258 genes. The RNAseq dataset consists of 60 498 probes. The mutation dataset is unified by gene identifiers across 40 543 genes. These datasets are then integrated to form a unified gene set. Afterwards, duplicate genes in each dataset are unified by taking mean values. All the datasets are filtered using Ensembl dataset. Integration with the KEGG database refines the data to 2117 common genes. Clinical dataset preprocessing involves dropping or imputing missing values, leading to M common samples for multi-omic analysis.

proximal promoter region, core promoter region, and downstream promoter region. After the unique genes in all the regions were integrated and padded, the final number of rows were 24 396 for all 5 files.

For UCSC genome variant datasets, values of duplicated genes were averaged as described in the method part to

represent the effect of variations of the same gene which reducing the number of rows from 40 543 to 40 490. With respect to UCSC RNA sequencing datasets, mean values were calculated for duplicate genes to ensure a single, representative value for each gene, the resulting number of rows from 60 498 to 49 432. The UCSC protein expression was also

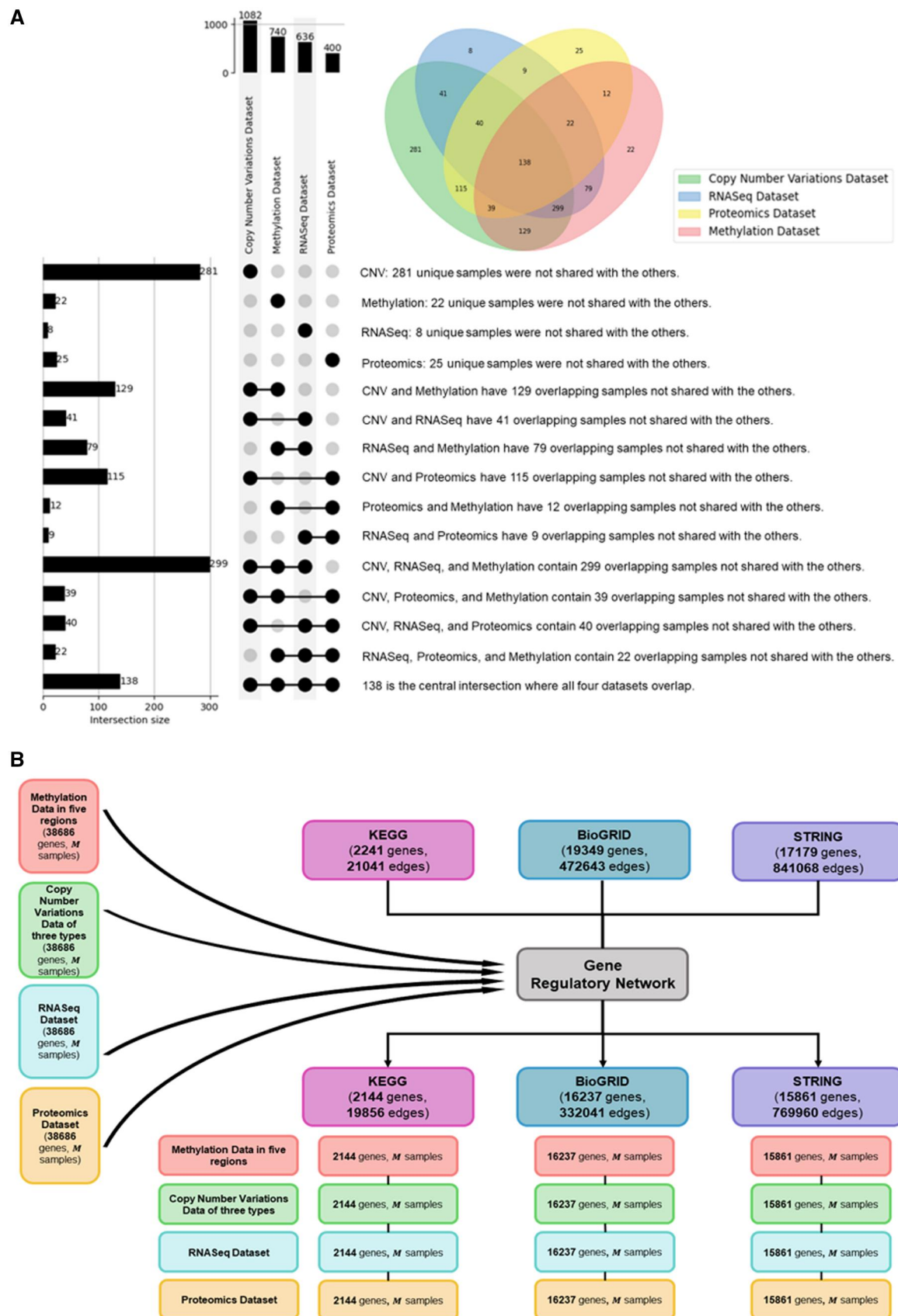


Figure 6. Details for intersection for samples gene regulatory network. (A) The intersection of omics data with clinical samples, finally M samples across all omics data. (B) The intersection of the multi-omic data with different options of gene regulatory network.

Table 4. ROSMAP Database resources.

Database	Description	Link
ROSMAP_arrayMethylation_imputed	Methylation data was generated on prefrontal cortex samples collected from 708 individuals using the Illumina HumanMethylation450 BeadChip	https://www.synapse.org/#!/Synapse:syn3168763
C2.median_polish_corrected_log2(Proteomics)	Data generated from isobaric TMT peptide labeling of ROSMAP brain tissues were submitted to the AD Knowledge Portal in two rounds. Round 1 (submitted in 2018) provides data from 400 individuals. Round 2 (submitted in 2022) provides data from an additional 210 individuals.	https://www.synapse.org/#!/Synapse:syn21266454
ROSMAP_RNAseq_FPKM_gene	Samples were extracted using Qiagen's miRNeasy mini kit (cat. no. 217004) and the RNase free DNase Set (cat. no. 79254), and quantified by Nanodrop and quality was evaluated by Agilent Bioanalyzer.	https://www.synapse.org/#!/Synapse:syn3505720
ROSMAP.CNV.Matrix(Mutation)	The TCGA Unified Ensemble “MC3” gene-level mutation dataset identifies somatic mutations in various cancers, marking nonsilent mutations (1) that alter protein sequences and wild type 0 for no mutations.	https://www.synapse.org/#!/Synapse:syn26263118
GEO GPL16304 Platform	Illumina HumanMethylation450 BeadChip [UBC enhanced annotation v1.0]	https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GPL16304
ROSMAP_clinical	Contains patient clinical features. A large amount of clinical and pathological data have been collected from individuals in the ROSMAP studies. The remainder of the clinical and pathological data may be accessed directly from the Rush Alzheimer's Disease Center.	https://www.synapse.org/#!/Synapse:syn3191087

Table 5. Dimensions of the ROSMAP multi-omics data before and after the data processing.

Database	Before process	After
Methylation dataset	(420 132 probes, 740 samples)	
Copy number variations dataset	(9902 chroms, 1127 samples)	(2144 genes in KEGG,
RNAseq dataset	(55 889 gene IDs, 640 samples)	16 237 genes in BioGRID,
Proteomics dataset	(8817 genes, 400 samples)	15 861 genes in STRING,
		<i>M</i> samples)
Clinical dataset	(16 features, 3584 samples)	(14 features, <i>M</i> samples)

checked with the duplicated gene names problem and found no duplication. Therefore, the number of rows (proteins) remains as the initial number as 258 (check Fig. 5 for details).

Furthermore, the reference datasets (clinical features, Ensembl dataset, gene regulatory network) were collected and integrated with processed multi-omics dataset for providing labels, gene annotation, and knowledge graph construction. Firstly, an examination of four clinical feature datasets sourced from the UCSC databases was conducted. The concatenation of these datasets yielded a composite dataset encompassing 6385 samples and 22 features, following the exclusion of any features exhibiting >30% missing values. In multi-omics dataset, the number of samples are 9664, 9104, 10 535, and 258 in methylation, mutation, RNASeq, and protein expression dataset. Afterward, the common samples/patients across multi-omics datasets and concatenated clinical dataset were intersected, leaving *M* common samples, where $M = 3592$ if all of the multi-omics datasets were utilized (see Fig. 6A and Table 3).

What's more, the Ensembl dataset was leveraged to annotate gene in multi-omics data, resulting in the dimensions for these files were methylation dataset (18 520 genes, *M*

samples), protein expression dataset (87 genes, *M* samples), RNASeq dataset (33 248 genes, *M* samples), and mutation dataset (30 275 genes, *M* samples). To obtain the union of genes from multi-omics data, all unique genes were collected across the datasets, resulting in a total of 34 366 genes. Then, the genes used to construct the knowledge graph were collected by intersecting genes in multi-omics datasets and gene regulatory network [KEGG (Ogata *et al.* 1999, Kanehisa and Goto 2000) (2241 genes, 21 041 edges), BioGRID (Stark *et al.* 2006, Oughtred *et al.* 2019) (19 349 genes, 472 643 edges), and STRING (Szklarczyk *et al.* 2019) (17 179 genes, 841 068 edges)], ending with the number of gene entities as 2117, 16 359, and 15 923 (see Fig. 6B).

3.2.3 Multi-omics datasets of Alzheimer's disease (AD)

Table 4 presents the multi-omics datasets and the corresponding dataset's download links. The mutation data, derived from “Whole genome sequencing-based copy number variations reveal novel pathways and targets in Alzheimer's disease,” were categorized into three distinct mutation types. For the RNA-seq data, null values will be replaced with 0 to ensure the data is meaningful and suitable for modeling

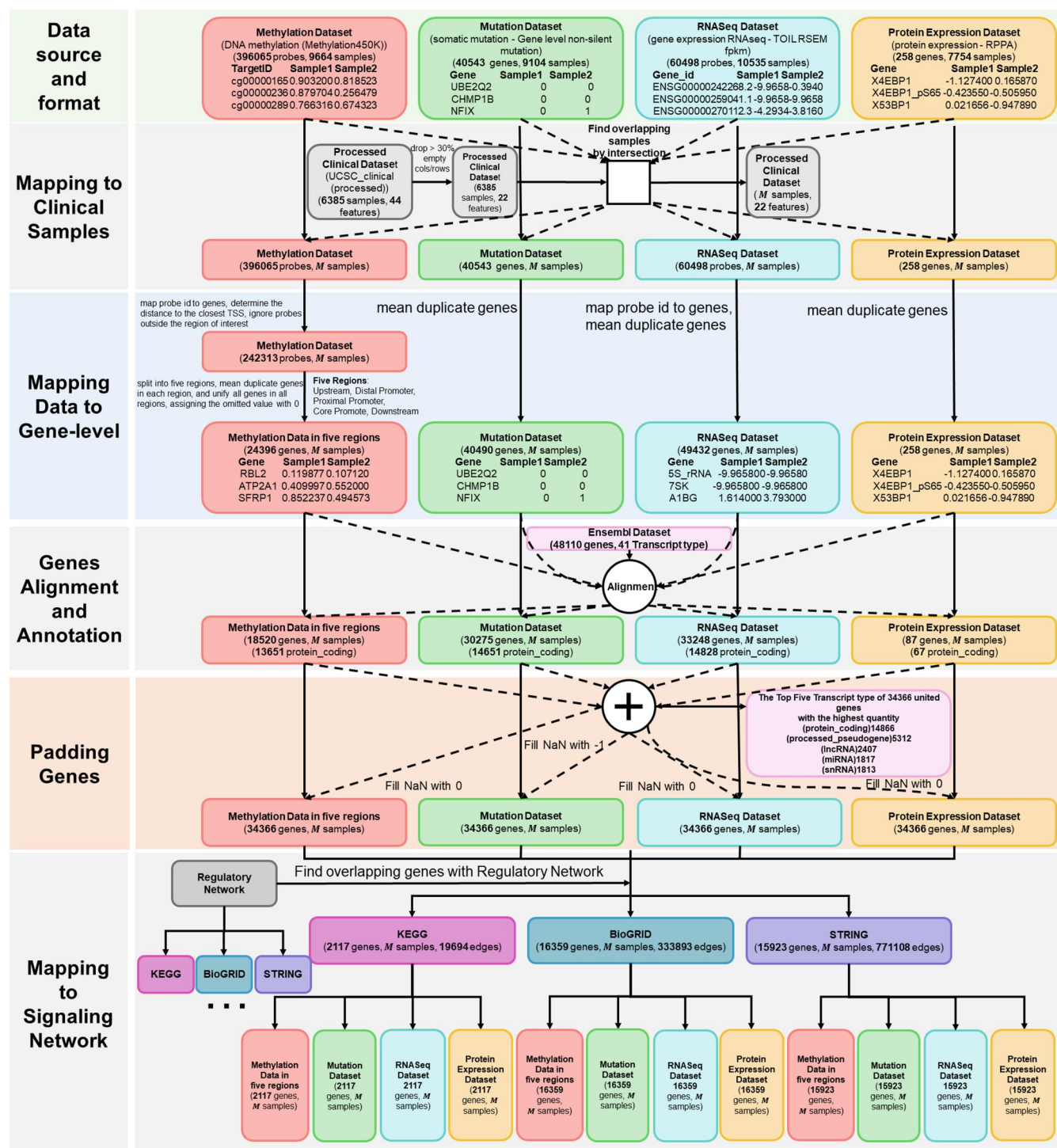


Figure 7. Flowchart of AD multi-omic data processing. The methylation dataset undergoes probe-to-gene mapping, with a focus on the -6 kb to $+3$ kb range around the TSS, followed by division into five genomic regions. The proteomics dataset includes 8817 genes. The RNAseq dataset consists of 55 889 probes. The copy number variations dataset undergoes chromosome position-to-gene mapping, followed by division into three types. These datasets are then integrated to form a unified gene set. Afterwards, duplicate genes in each dataset are unified by taking mean values. All the datasets are filtered using Ensembl dataset. Integration with the KEGG database refines the data to 2144 common genes. Clinical dataset preprocessing involves dropping or imputing missing values, leading to M common samples for multi-omic analysis.

purposes. A crucial step in the preprocessing stage was the unification of genes across each dataset. Once all processes were completed, all datasets were collectively processed. The detailed pre-analysis pipeline is described as follows.

As described in the method aforementioned, raw multi-omics datasets in ROSMAP are collected and loaded into the pandas dataframe (check Table 5 for rows and columns information). The methylation dataset encompasses 420 132

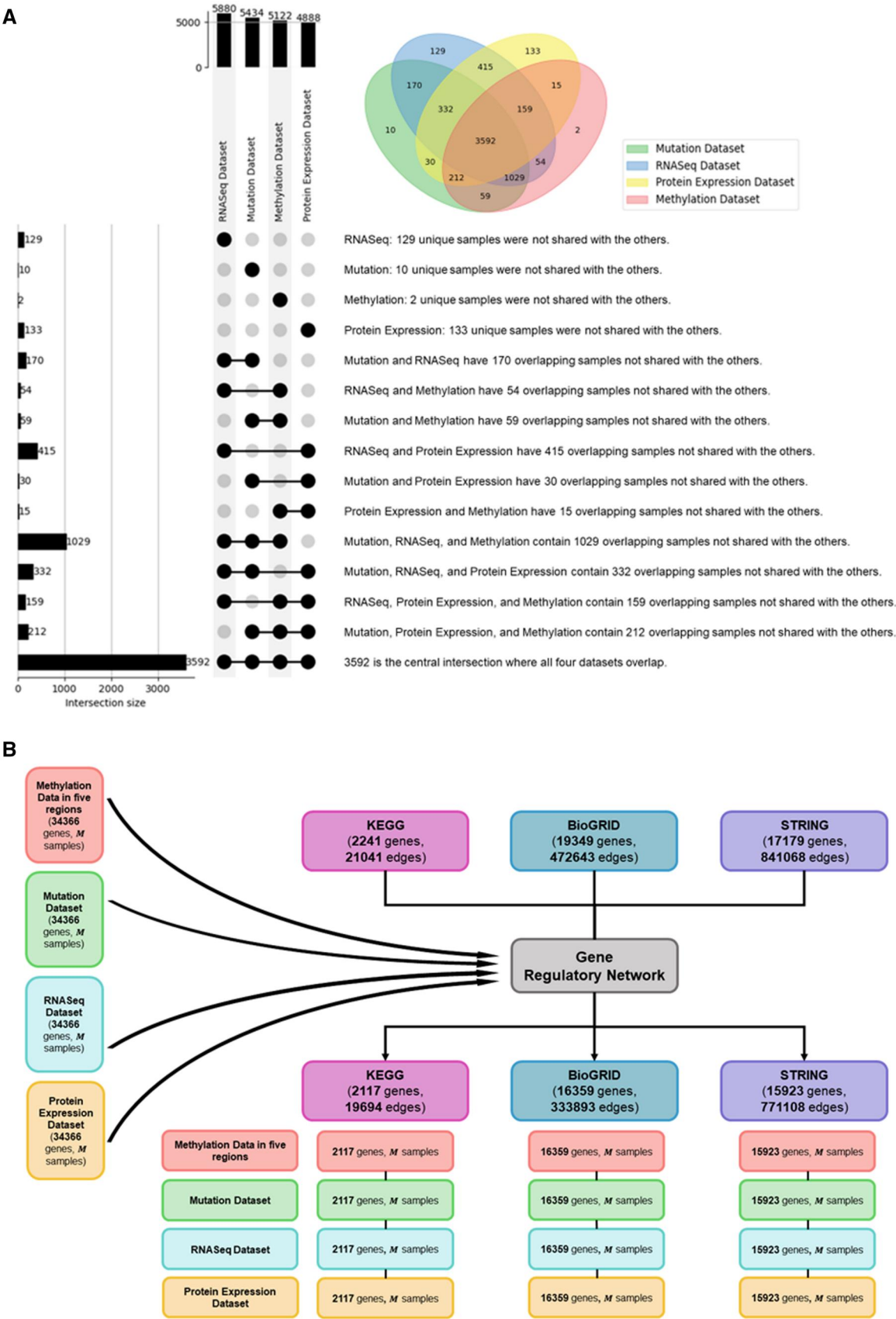


Figure 8. Details for intersection for samples gene regulatory network. (A)The intersection of omics data with clinical samples, finally M samples across all omics data. (B) The intersection of the multi-omic data with different options of gene regulatory network.

probes, targeting the CpG sites of genes. Since the regions of interest around TSS were predefined as mentioned in Fig. 2, rows with TSS distances that fell outside these areas were excluded and each CpG probe was assigned with corresponding TSS closest gene name reducing the number of rows from 420 132 to 249 506. By splitting the methylation files into 5 files according to predefined promoter regions, the number of rows (genes) will be 15 812, 20 577, 18 238, 9871, and 7128 for the upstream region, distal promoter region, proximal promoter region, core promoter region, and downstream promoter region.

After the unique genes in all the regions were integrated and padded, the final number of rows were 23 280 for all 5 files.

With respect to Copy Number Variations dataset, data were aggregated with summation for duplicate genes to capture the effect of variations across multiple instances of the same gene. To split and unify data, genome variants data were categorized based on the nature of the variation—deletions (DEL), duplications (DUP), and multiple copy number variations (mCNV). And there were three final aggregated files for “DEL,” “DUP,” and “mCNV” with a number of rows with 2483, 1096, and 3726. By identifying and adding all unique gene names to each file, the three files had the same format, generating 6423 unique genes for all 3 types of variation. For any duplicate genes in ROSMAP RNA sequencing data, mean values were calculated to ensure a single, representative value for each gene, resulting number of rows from 55 889 to 37 898. The proteomics in the ROSMAP dataset was also examined for the problem of duplicated gene names, and gene uniqueness was ensured by merging duplicates, reducing the gene count from 8817 to 8252 (check Fig. 7 for details).

Furthermore, the reference datasets (clinical features, Ensembl dataset, gene regulatory network) were collected and integrated with processed multi-omics dataset for providing labels, gene annotation, and knowledge graph construction. The common samples/patients across multi-omics datasets and concatenated clinical dataset were intersected, leaving M common samples, where $M = 138$ if all of the multi-omics datasets were utilized (see Fig. 8A and Table 5). Additionally, the Ensembl dataset was utilized to annotate genes in multi-omics data, which led to the following dimensions: methylation dataset (17 606 genes, M samples), Copy Number Variations dataset (6419 genes, M samples), RNASeq dataset (37 898 genes, M samples), and proteomics dataset (7817 genes, M samples). The union of genes from the multi-omics datasets was achieved by aggregating all unique genes across the datasets, yielding a total of 38 686 genes. Subsequently, genes for constructing the knowledge graph were selected through the intersection of multi-omics datasets with gene regulatory networks—specifically, [KEGG (Ogata *et al.* 1999, Kanehisa and Goto 2000) (2241 genes, 21 041 edges), BioGRID (Stark *et al.* 2006, Oughtred *et al.* 2019) (19 349 genes, 472 643 edges), and STRING (Szklarczyk *et al.* 2019) (17 179 genes, 841 068 edges)], resulting the number of gene entities as 2144, 16 237, and 15 861 (see Fig. 8B).

4 Discussion

Our study introduced mosGraphGen, a novel computational tool that enables the generation of multi-omics signaling graphs for each sample by mapping the multi-omics data

onto a biologically relevant multi-level signaling network. With mosGraphGen, AI model developers can seamlessly apply and assess their models using these multi-omics signaling graphs. We conducted evaluations of mosGraphGen using multi-omics datasets from both cancer and Alzheimer’s disease (AD) samples. Furthermore, we have made the code, examples, and tutorials for mosGraphGen open-source and publicly accessible for the research community via GitHub: <https://github.com/FuhaiLiAiLab/mosGraphGen>. The following sections provide the details of the methodology and results. Users can build the mosGraphs using the two widely used multi-omics datasets following the instructions for their own down-stream analysis.

Given the exploratory nature of this study, it is important to acknowledge the inherent limitations that present opportunities for in-depth exploration in subsequent research endeavors. For instance, there is potential for enhancing the biological significance of the background signaling network to improve the biological relevance of multi-omics data integration. It is also important to biologically or experimentally validate these mosGraphs in specific studies, which can further indicate the reliability and utility, and further improve the model. Furthermore, we aspire to develop an interactive tool that allows users to generate mosGraphs without the need for coding, particularly tailored for well-known and widely used multi-omics datasets and data platforms. There are a few potential future improvements and extensions of mosGraphGen model. For example, the nucleic acid sequence can be included to indicate the specific and fine-level genetic mutations and epigenetic changes. It will increase the computational cost dramatically. Also, the metabolomic data is critical to understand disease pathogenesis. Thus, it is also important to integrate metabolomic data using the known-protein metabolites interactions or their co-expressions to identify the synergistic effects of the multi-omic and metabolomic features. Moreover, the single-cell omic datasets have been generated to understand disease molecular mechanisms at single-cell level and potential cell-cell interactions. Thus, it is interesting and important to expand the mosGraphs for single-cell omics data. However, it’s important to note that dealing with a substantially larger number of mosGraphs for single-cell multi-omics data, given the presence of tens of thousands of scs in individual studies, poses a challenge in terms of data storage and analysis. In addition to the numeric multi-omic data, the known knowledge or descriptive/text information of individual proteins, like the annotation and description of the functions of individual proteins, can be added. The large language models (LLMs) can be used to convert the description/text information into numeric feature to be combined with the numeric omic features for the target and signaling pathway inference/mining from the mosGraphs. Therefore, novel graph AI models that can analyze the large-scale mosGraphs and identify the essential disease-associated targets and signaling pathways are needed.

Acknowledgements

H.Z., D.C., Z.C., and X.Z. contributed to data curation, methodology, formal analysis, visualization, and writing the original draft. H.Z. and D.C. were responsible for reviewing the manuscript. F.L. contributed to conceptualization, writing, and manuscript review. Y.C., C.S., C.C., P.P., G.L., and

M.P. contributed to the methodology and investigation. All authors read and approved the final manuscript.

Conflict of interest

The authors declare no conflict of interest.

Funding

This work was partially supported by R56AG065352, 1R21AG078799-01A1, 1RM1NS132962-01, 1R01LM013902-01A1.

Data availability

Datasets used in this research were collected from public websites. Please check Table 2 and Table 4 for details.

References

- Allen M, Carrasquillo MM, Funk C *et al.* Human whole genome genotype and transcriptome data for Alzheimer's and other neurodegenerative diseases. *Sci Data* 2016;3:160089. <https://doi.org/10.1038/sdata.2016.89>
- Barretina J, Caponigro G, Stransky N *et al.* The cancer cell line encyclopedia enables predictive modelling of anticancer drug sensitivity. *Nature* 2012;483:603–7. <https://doi.org/10.1038/nature11003>
- Baysoy A, Bai Z, Satija R *et al.* The technological landscape and applications of single-cell multi-omics. *Nat Rev Mol Cell Biol* 2023;24:695–713. <https://doi.org/10.1038/s41580-023-00615-w>
- Bennett DA, Buchman AS, Boyle PA *et al.* Religious orders study and rush memory and aging project. *J Alzheimers Dis* 2018;64:S161–89.
- Butler JEF, Kadonaga JT. The RNA polymerase II core promoter: a key component in the regulation of gene expression. *Genes Dev* 2002;16:2583–92. <https://doi.org/10.1101/gad.1026202>
- De Jager PL, Ma Y, McCabe C *et al.* A multi-omic atlas of the human frontal cortex for aging and Alzheimer's disease research. *Sci Data* 2018;5:1–13.
- Deelen J, Evans DS, Arking DE *et al.* A meta-analysis of genome-wide association studies identifies multiple longevity genes. *Nat Commun* 2019;10:3669. <https://doi.org/10.1038/s41467-019-11558-2>
- Dong Z, Zhang H, Chen Y *et al.* Interpreting the mechanism of synergism for drug combinations using attention-based hierarchical graph pooling. *Cancers* 2023;15:4210. <https://doi.org/10.3390/cancers15174210>
- Duttke SHC, Lacadie SA, Ibrahim MM *et al.* Human promoters are intrinsically directional. *Mol Cell* 2015;57:674–84. <https://doi.org/10.1016/j.molcel.2014.12.029>
- Gao H, Zhang B, Liu L *et al.* A universal framework for single-cell multi-omics data integration with graph convolutional networks. *Brief Bioinform* 2023;24:bbad081.
- Goldman MJ, Craft B, Hastie M *et al.* Visualizing and interpreting cancer genomics data via the Xena platform. *Nat Biotechnol* 2020;38:675–8. <https://doi.org/10.1038/s41587-020-0546-8>
- Greenwood AK, Montgomery KS, Kauer N *et al.* The AD knowledge portal: a repository for multi-omic data on Alzheimer's disease and aging. *Curr Protoc Hum Genet* 2020;108:e105. <https://doi.org/10.1002/cphg.105>
- Grossman RL, Heath AP, Ferretti V *et al.* Toward a shared vision for cancer genomic data. *N Engl J Med* 2016;375:1109–12. <https://doi.org/10.1056/nejmp1607591>
- Kanehisa MGS, Goto S. KEGG: Kyoto encyclopedia of genes and genomes. *Nucleic Acids Res* 2000;28:27–30. <https://doi.org/10.1093/nar/28.1.27>
- Kuenzi BM, Park J, Fong SH *et al.* Predicting drug response and synergy using a deep learning model of human cancer cells. *Cancer Cell* 2020;38:672–84.e6. <https://doi.org/10.1016/j.ccell.2020.09.014>
- Lee J, Hyeon DY, Hwang D. Single-cell multiomics: technologies and data analysis methods. *Exp Mol Med* 2020;52:1428–42.
- 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. *Front Aging Neurosci* 2022a;14:935279.
- 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* 2022b;13:806842.
- Ma J, Yu MK, Fong S *et al.* Using deep learning to model the hierarchical structure and function of a cell. *Nat Methods* 2018;15:290–8. <https://doi.org/10.1038/nmeth.4627>
- Neff RA, Wang M, Vatansever S *et al.* Molecular subtyping of Alzheimer's disease using RNA sequencing data reveals novel mechanisms and targets. *Sci Adv* 2021;7:eabb5398. <https://www.sciencedirect.com/science/article/pii/S2352396421000000>
- Ogata H, Goto S, Sato K *et al.* KEGG: Kyoto encyclopedia of genes and genomes. *Nucleic Acids Res* 1999;27:29–34. <https://doi.org/10.1093/nar/27.1.29>
- Oughtred R, Stark C, Breitkreutz B-J *et al.* The BioGRID interaction database: 2019 update. *Nucleic Acids Res* 2019;47:D529–41.
- 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. *J Gerontol A Biol Sci Med Sci* 2022;77:2272–80. <https://doi.org/10.1093/gerona/glac065>
- Rajadhyaksha N, Chitkara A. Graph contrastive learning for multi-omics data. arXiv, arXiv:2301.02242, 2023, preprint: not peer reviewed.
- Saintenac C, Jiang D, Akhunov ED. Targeted analysis of nucleotide and copy number variation by exon capture in allotetraploid wheat genome. *Genome Biol* 2011;12:R88. <https://doi.org/10.1186/gb-2011-12-9-r88>
- Sanchez-Vega F, Mina M, Armenia J *et al.* Cancer Genome Atlas Research Network. Oncogenic signaling pathways in the cancer genome atlas. *Cell* 2018;173:321–37.e10. <https://doi.org/10.1016/j.cell.2018.03.035>
- Saxonov S, Berg P, Brutlag DL. A genome-wide analysis of CpG dinucleotides in the human genome distinguishes two distinct classes of promoters. 2005. www.pnas.org/doi/10.1073/pnas.0510310103
- Sedgewick AJ, Benz SC, Rabizadeh S *et al.* Learning subgroup-specific regulatory interactions and regulator independence with PARADIGM. *Bioinformatics* 2013;29:i62–70.
- Slenter DN, Kutmon M, Hanspers K *et al.* WikiPathways: a multifaceted pathway database bridging metabolomics to other omics research. *Nucleic Acids Res* 2018;46:D661–7.
- Stark C, Breitkreutz B-J, Reguly T *et al.* BioGRID: a general repository for interaction datasets. *Nucleic Acids Res* 2006;34:D535–9.
- Subramanian I, Verma S, Kumar S *et al.* Multi-omics data integration, interpretation, and its application. *Bioinform Biol Insights* 2020;14:1177932219899051.
- Szklarczyk D, Franceschini A, Wyder S *et al.* STRING v10: protein–protein interaction networks, integrated over the tree of life. *Nucleic Acids Res* 2015;43:D447–52.
- Szklarczyk D, Gable AL, Lyon D *et al.* STRING v11: protein–protein association networks with increased coverage, supporting functional discovery in genome-wide experimental datasets. *Nucleic Acids Res* 2019;47:D607–13.
- Vanaja A, Yella VR. Delineation of the DNA structural features of eukaryotic core promoter classes. *ACS Omega* 2022;7:5657–69. <https://doi.org/10.1021/acsomega.1c04603>
- van den Berg R, Kipf TN, Welling M. 'Graph convolutional matrix completion'. KDD'18 Deep Learning Day, 2018. <https://arxiv.org/abs/1706.02263>
- 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;12:3445.
- Yu MK, Ma J, Fisher J *et al.* Visible machine learning for biomedicine. *Cell* 2018;173:1562–5. <https://doi.org/10.1016/j.cell.2018.05.056>
- Zhang H, Chen Y, Payne P *et al.* Using DeepSignalingFlow to mine signaling flows interpreting mechanism of synergy of cocktails. *NPJ Syst Biol Appl* 2024;10:92.

© The Author(s) 2024. Published by Oxford University Press.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

Bioinformatics Advances, 2024, 00, 1–14

<https://doi.org/10.1093/bioadv/vbae151>

Original Article