



OPEN ACCESS

EDITED BY

An Zhu,
Fujian Medical University, China

REVIEWED BY

Lin Zhang,
China University of Mining and Technology,
China
Songyao Zhang,
National University of Singapore, Singapore

*CORRESPONDENCE

Fei Liu,
✉ bwlif@163.com
Lian Liu,
✉ lian.l@snnu.edu.cn

RECEIVED 09 April 2025

ACCEPTED 12 May 2025

PUBLISHED 27 May 2025

CITATION

Bai L, Liu F, Wang Y, Su J and Liu L (2025) MultiV_Nm: a prediction method for 2'-O-methylation sites based on multi-view features.
Front. Genet. 16:1608490.
doi: 10.3389/fgene.2025.1608490

COPYRIGHT

© 2025 Bai, Liu, Wang, Su and Liu. This is an open-access article distributed under the terms of the [Creative Commons Attribution License \(CC BY\)](#). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

MultiV_Nm: a prediction method for 2'-O-methylation sites based on multi-view features

Lei Bai¹, Fei Liu^{1*}, Yile Wang¹, Junle Su¹ and Lian Liu^{2*}

¹School of Physics and Opto-Electronic Technology, Baoji University of Arts and Sciences, Baoji, China,

²School of Computer Sciences, Shannxi Normal University, Xi'an, China

As a crucial class of chemical modifications, 2'-O-methylation modification (abbreviated as Nm) is widely distributed in various organisms and plays a very important role in normal cellular physiological activities and the occurrence and development of diseases. Accurate prediction of Nm modification sites can provide important references for the diagnosis and treatment of diseases, as well as identifying for potential drug targets. Aiming at the current problems of unstable performance caused by the use of single features and the need to improve the prediction accuracy of Nm modification sites, this paper proposes MultiV_Nm, a prediction method for Nm sites based on multi-view features. MultiV_Nm extracts the features of Nm sites from multiple dimensions, including sequence features, chemical characteristics, and secondary structure features. By integrating the powerful local feature extraction ability of convolutional neural networks, the ability of graph attention networks to capture global structural information, and the efficient interaction advantage of cross-attention mechanisms for different features, it deeply explores and integrates multi-view features, and finally realizes the prediction of Nm modification sites. The results of cross-validation and independent tests show that this method exhibits significant advantages in key evaluation indicators such as precision, recall, and accuracy, and can effectively improve Nm sites prediction performance. The proposal of MultiV_Nm not only provides a powerful tool for the study of Nm modification but also offers new ideas for predicting other RNA modification sites.

KEYWORDS

2'-O-methylation sites, multi-view, convolutional neural networks, graph attention network, cross attention mechanism

1 Introduction

In recent years, with the in-depth research in the field of epitranscriptomics, RNA modifications, as a crucial epigenetic mechanism for gene expression regulation, have attracted extensive attention in the field of biomedical research (Boccalletto et al., 2018). Different from DNA and histone modifications, RNA modifications regulate various aspects of RNA metabolism through dynamic and reversible chemical modification methods, have a significant impact on biological processes such as individual development and disease occurrence (Roundtree et al., 2017), and are closely related to the occurrence and development of a variety of diseases (Li et al., 2018; Cui et al., 2022). Currently, more than 200 different types of RNA chemical modifications have been identified in eukaryotes (Boccalletto et al., 2022). Among them, 2'-O-methylation (abbreviated as Nm) is an extremely important and widely existing type of RNA

modification. Catalyzed by 2'-O-methyltransferases, it adds a methyl group to the 2'-hydroxyl group of RNA (Nicolai et al., 2016). Nm exists in the 2'-hydroxyl ribose moieties of all four ribonucleosides (Xuan et al., 2018), namely, 2'-O-methylcytidine (Cm), 2'-O-methyladenosine (Am), 2'-O-methylguanosine (Gm), and 2'-O-methyluridine (Um). Widely present in various RNA molecules (Li et al., 2024), this modification plays a key role in maintaining normal physiological functions in organisms (Yu et al., 2004). The distribution of Nm is extremely extensive, and it can be found in rRNA, mRNA, tRNA, snRNA, piRNA as well as human viruses (Wu et al., 2024). In rRNA, the enzyme fibrillarin (FBL) can catalyze the Nm reaction. In the breast cancer cell model, the mutation of the tumor suppressor gene TP53 will increase the expression of FBL, which in turn leads to an elevated level of Nm modification in rRNA and abnormal translation of the internal ribosome entry site (IRES) oncogenes (Marcel et al., 2013). In mRNA, Nm modifications occur both at the 5' cap and internal sites. The modification of the 5' cap can protect the mRNA and regulate immune recognition; the modifications at internal sites can affect the translation efficiency and mRNA stability, and are also associated with viral infections (Picard-Jean et al., 2018). The Nm modification can stabilize the L-shaped tertiary structure of tRNA, enhance its thermal stability, and contribute to its correct folding. Moreover, it will also affect the recognition process of the codon by the tRNA anticodon during translation (Agris et al., 2017). The dysregulation of Nm modification in snRNA may disrupt the splicing process, resulting in the generation of erroneous mRNA and protein sequences, and may potentially trigger diseases such as cancer and splicing-related genetic disorders (Blijlevens et al., 2021). The study by Lim et al. has demonstrated that the protein HENMT1 is a key regulator of Nm modification in mammalian piRNA (Lim et al., 2015). Moreover, the Nm modification of the genomes of RNA viruses such as human immunodeficiency virus type 1 (HIV-1) and West Nile virus (WNV) may help them evade the host's innate immune response (Daffis et al., 2010; Ringear et al., 2019), providing potential targets for antiviral treatment. In addition, a large number of studies have continuously shown that RNA Nm modification is closely related to a variety of human diseases, such as hepatocellular carcinoma and lung adenocarcinoma (Song et al., 2023).

To deeply explore the functional mechanism of Nm modification, researchers have developed a series of biological experimental techniques (Yu et al., 1997; Maden et al., 1995; Zhang et al., 2023). However, these traditional techniques generally suffer from the problems of long time consumption and high cost, and it is difficult to meet the needs of biological research for efficient and rapid detection. With the rapid development of sequencing technologies, nucleotide sequence data have shown explosive growth, which has opened up a new path for computational prediction methods. Predicting Nm modification sites through computation can effectively make up for the deficiencies of experimental techniques and provide strong support for relevant research.

Currently, the computational tools for predicting RNA Nm modification sites are relatively limited. Chen et al. (2016) constructed the first computational tool for identifying Nm modification sites by using support vector machine (SVM) for classification based on the encoding methods of nucleotide

chemical properties and nucleotide composition features. However, this model was only constructed based on human data, and its prediction performance in other species has not been fully demonstrated. Qiu et al. (2017) incorporated the sequence coupling effect into the General Pseudo k-tuple Nucleotide Composition (General PseKNC) and used a variety of machine learning algorithms to construct an ensemble classifier. By combining and optimizing different algorithms, the prediction accuracy and stability were improved. In 2018, Yang et al. (2018) developed a sequence-based predictor iRNA-2OM for humans. By fusing chemical properties, nucleotide composition and PseKNC features, and combining feature selection methods with incremental feature selection (IFS), they obtained the optimal feature set and then constructed the prediction model. Zhou et al. (2019) developed the predictor NmSEER2.0 for Nm modification sites in the genomes of human HeLa and HEK293 T cells. This tool is based on random forest (RF) and multiple encoding schemes, and its AUC value reaches 0.862, showing good prediction performance. Ao et al. (2022) constructed the predictor NmRF based on the optimal mixed features and random forest classifier. By fusing nucleotide chemical properties, binary features and dinucleotide position-specific features, and then using a two-step strategy of combining the light gradient boosting algorithm with IFS feature selection, NmRF obtained the optimal feature set.

In addition, deep learning has gradually been applied to this field. The Deep-2'-O-Me method proposed by Mostavi et al. (2018) uses the dna2vec embedding method improved based on the word2vec model to learn the complex feature representations of pre-mRNA sequences, and fine-tunes them with the help of a convolutional neural network (CNN). In the test scenarios of both balanced and imbalanced datasets, the AUC and AUPRC scores both reach 0.9, significantly outperforming existing algorithms. iRNA-PseKNC (Tahir et al., 2019) uses PseKNC to extract the features of RNA sequences, and utilizes the feature learning ability of convolutional neural networks to automatically extract deep-level features and explore the complex relationships between RNA sequences and Nm sites. DeepOMe (Li et al., 2021) combines CNN and bidirectional long short-term memory networks (BLSTM), which enables it to accurately predict the Nm sites in the human transcriptome. Pichot et al. (Pichot et al., 2022) utilized the RiboMethSeq dataset and employed the random forest algorithm to construct a predictive model for analyzing Nm sites in RNA. This model was trained on a large number of human rRNA datasets with known modification profiles, and the modification profiles of other eukaryotic rRNAs (*Saccharomyces cerevisiae* and *Arabidopsis thaliana*) determined through experiments were used to evaluate the performance of the predictive model. For each type of Nm methylation, i2OM (Yang et al., 2023) combines one-way analysis of variance with mutual information to rank the sequence features, to obtain the optimal feature subset. Subsequently, four predictors based on eXtreme Gradient Boosting (XGBoost) or SVM are used to identify four types of Nm sites. BERT2OME (Soylu and Sefer, 2023) combines the BERT-based model with CNN to infer the relationship between the modification sites and the RNA sequence content. The results show that BERT2OME reduces the time consumed in biological experiments, and outperforms existing methods in terms of multiple metrics across different datasets and species. A

TABLE 1 Single-base resolution datasets in Nm prediction.

Id	Cell	Note	Technique	Source
1	HeLa	Control	Nm-seq	Dai et al. (2017)
2	HEK293 T	Control	Nm-seq	Dai et al. (2017)

large number of cutting-edge studies have shown that Nm modification is widely involved in key biological processes such as RNA splicing, transportation, and stability regulation. Accurately identifying Nm methylation modification sites helps deeply understand the pathogenesis of diseases and provides an important basis for developing new diagnostic and treatment strategies. However, current prediction methods for Nm methylation sites have obvious deficiencies. Most existing models rely only on single features, making it hard to comprehensively capture information in RNA sequences and structures, leaving significant room for improving the models' prediction accuracy and stability.

To overcome this technical hurdle, we propose MultiV_Nm, an innovative prediction framework for Nm methylation sites using multi-view features. It extracts nucleotide sequence features through one-hot encoding, explores chemical properties and obtains RNA secondary structural features. The model combines convolutional neural networks for local feature extraction, graph attention networks for global relationship modeling, and a cross-attention mechanism for feature interaction. This integration enables in-depth understanding of multi-view features, significantly enhancing the prediction accuracy of Nm methylation sites.

2 Materials and methods

2.1 Datasets construction

To predict the Nm methylation modification sites, we utilized the Nm-seq technology to collect the information of Nm methylation modification sites at single-base resolution from two types of cells, namely, HeLa and HEK293T cells (see Table 1). During the research process, the Nm methylation sites detected in these two types of cells were defined as positive Nm methylation modification sites. To construct the negative sample set, we randomly selected an equal number of sites as the positive sites from the regions containing the positive samples. At the same time, we excluded the sites located in the ambiguous regions that could be mapped to multiple genes to ensure the accuracy and reliability of the data.

According to the statistics, a total of 7,193 positive Nm methylation sites were finally collected. Among them, there were 1,591 Am type sites, accounting for 22.1% of all Nm sites; 1,471 Gm type sites, accounting for 20.5%; 1878 Cm type sites, accounting for 26.1%; and 2,253 Um type sites, accounting for 31.3%. To ensure the balance between positive and negative samples, the number of negative samples is the same as the positive. Subsequently, 3,696 samples were extracted from the collected samples as the test samples, including 1,848 positive samples and 1,848 negative samples. According to the proportion of each type of Nm in the total Nm, the numbers of Am, Gm, Cm, and Um in the test samples were 352, 300, 490, and 706, respectively, and the number of negative samples was the same as the positive. The remaining samples were

used as the training samples, totaling 10,690, including 5,345 positive samples and 5,345 negative samples.

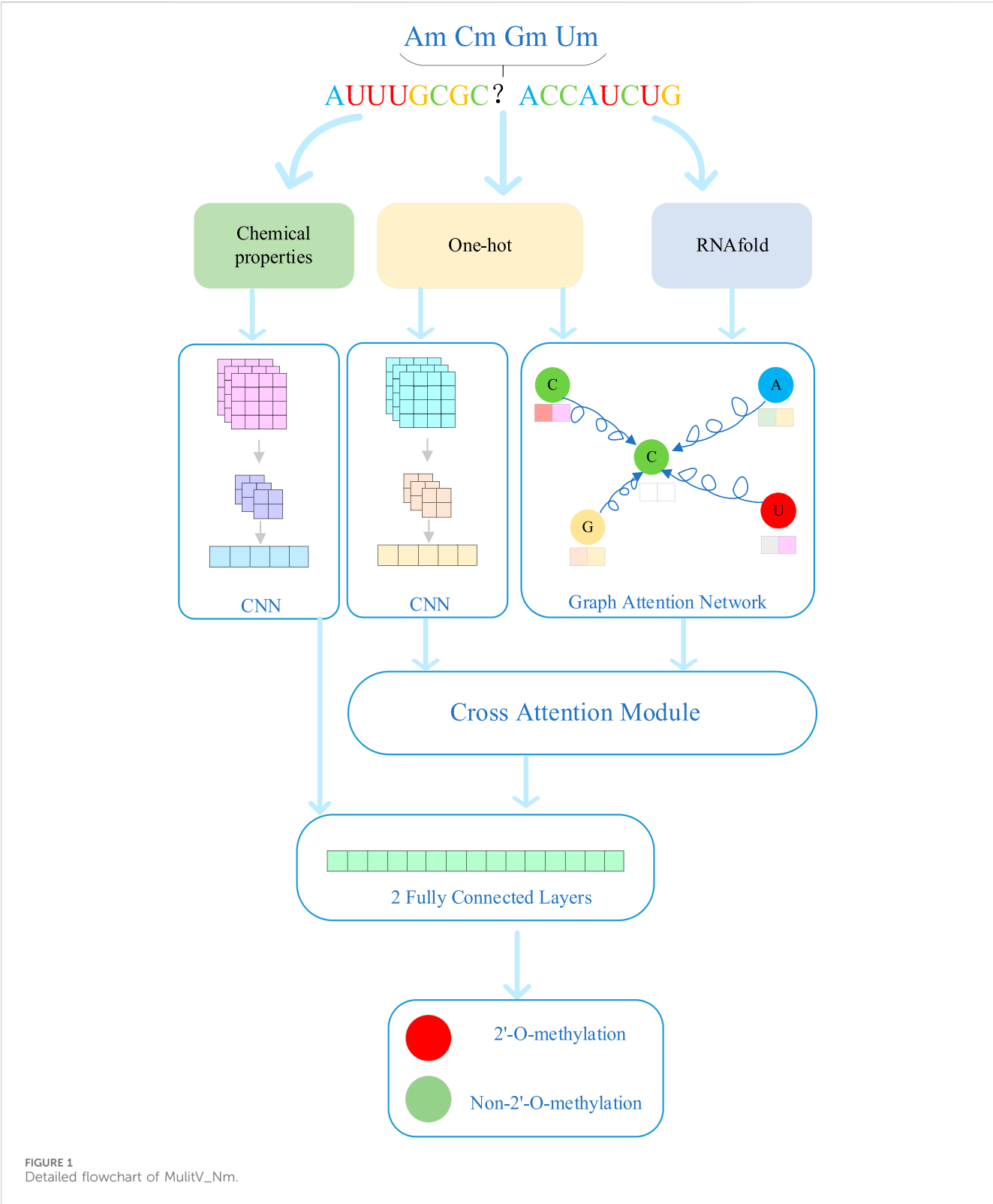
To obtain the sequence data for model training and testing, taking the Nm methylation sites as the center, we extended 25 base pairs (bp) upstream and downstream respectively, and finally obtained the Nm methylation modification site sequences with a length of 51 bp, laying a solid data foundation for the subsequent prediction analysis.

2.2 Methods

2.2.1 Overall model architecture

The overall model architecture of MultiV_Nm is shown in Figure 1. First, for the known RNA sequences, the model extracts the features from sequence and chemical properties. For sequence features, one-hot encoding is used for feature extraction, which completely preserves the genetic and regulatory information contained in the arrangement order of nucleotides. At the same time, a quantitative analysis is carried out on the chemical property features of RNA molecules to explore their potential associations with modification sites. In addition, through the RNAfold software, the secondary structure features of RNA are analyzed to obtain the local folding information formed by base pairing of the molecules.

The prediction process of MultiV_Nm involves four modules: the convolutional neural network (CNN), the graph attention network (GAT), the cross-attention mechanism, and the fully connected layers. The CNN is mainly used to extract the deep features of sequences and chemical properties. The GAT extracts the deep features of the secondary structure through the spatial relationships and connection information between nodes. The cross-attention module fuses the sequence features and secondary structure features to achieve feature complementarity. The fully connected layers obtain the final prediction results by integrating the chemical property features and the complementary features. First, since the CNN has a powerful ability of automatic learning in local feature extraction, MultiV_Nm combines CNN with the pooling layer to deeply mine the sequence features and chemical property features extracted in the early stage. It captures the deep features hidden in the data and enhances the expression ability of sequence features and chemical property features. Second, the GAT provides strong support for the analysis of secondary structure features. The RNA secondary structure is composed of many nodes. By introducing the attention mechanism, GAT can make full use of the information of nodes and edges, accurately capture the interactions of nucleotides in the spatial structure, and explore the deep features of the secondary structure. Third, the sequence features mainly carry genetic information and regulatory information of biological processes, while the secondary structure features intuitively show the local folding morphology of RNA molecules. To give full play to the advantages of both, MultiV_Nm introduces a cross-attention module to fuse the deep sequence features and structural features. This module can automatically learn the relationships between the two types of features, adaptively adjust the fusion weights, and achieve efficient integration of features. Finally, the deep chemical property features are concatenated with the features fused by the cross-attention module. Through two fully connected layers, the integrated features are further analyzed and processed, and finally, the prediction results of Nm methylation modification sites are output.



2.2.2 Feature representation

2.2.2.1 Sequence feature representation

One-hot encoding is a technique for converting categorical variables into vector representations. For a categorical variable with n different categories, One-hot encoding represents each category as a vector of length n , in which only one element is

1 and the rest are all 0. The RNA sequence is composed of adenine (A), cytosine (C), guanine (G), and uracil (U), and its character set is {A, C, G, U}. We assign a unique integer index to each character in the character set and create a dictionary to map characters to indices. The dictionary is created as follows: {A: 0, C: 1, G: 2, U: 3}. In this dictionary, each base corresponds to a

unique integer index, which facilitates subsequent encoding operations.

For each character, its corresponding index is obtained according to the encoding dictionary. Then, the value is set to 1 at the corresponding position in the feature vector. Suppose the index corresponding to the current character is i , then the i th element in the feature vector is set to 1. Therefore, A is encoded as [1, 0, 0, 0], C is encoded as [0, 1, 0, 0], G is encoded as [0, 0, 1, 0], and U is encoded as [0, 0, 0, 1]. For example, a sequence $seq = [AACUG]$ can be encoded as the matrix shown in Equation 1. For a sequence with a length of 51 bp, it is encoded as a 51×4 matrix through one-hot encoding.

$$seq = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix} \quad (1)$$

2.2.2.2 Chemical property features representation

Each nucleotide in RNA can be represented by three features according to its different chemical properties (Liu et al., 2020). C and U have only one ring structure, while A and G have two rings; both A and C contain an amino group, while both G and U contain a keto group; when forming the secondary structure, the hydrogen bonds between G and C are relatively strong, while the hydrogen bonds between A and U are relatively weak. Based on these three features, a nucleotide can be represented by a three-dimensional vector $S = (x_i, y_i, z_i)$, such as Equation 2):

$$x = \begin{cases} 1, s \in \{A, G\} \\ 0, s \in \{C, U\} \end{cases}, y = \begin{cases} 1, s \in \{A, C\} \\ 0, s \in \{G, U\} \end{cases}, z = \begin{cases} 1, s \in \{A, U\} \\ 0, s \in \{C, G\} \end{cases} \quad (2)$$

Therefore, A, C, G, and U can be encoded as [1, 1, 1], [0, 1, 0], [1, 0, 0], [0, 0, 1], and [0, 0, 1], respectively.

2.2.2.3 Secondary structure features representation

To extract the secondary structure features of RNA sequences, RNAfold (Lorenz et al., 2011) in the ViennaRNA package (<http://rna.tbi.univie.ac.at/cgi-bin/RNAWebSuite/RNAfold.cgi>) was used to predict the RNA secondary structure. According to the given RNA sequence, RNAfold can predict the possible secondary structure of RNA by calculating thermodynamic parameters and return the results in the form of dot-bracket. In order to show the relationship between bases in the RNA sequence, we constructed a base-base relationship matrix. In the dot-bracket, a “dot” represents an unpaired base, and it is set to 0 in the base-base relationship matrix. The left parenthesis “(“ and the right parenthesis”)” are used to represent paired bases. The left and right parentheses appear in pairs. The left parenthesis is placed at the position of one of the paired bases, and the right parenthesis is placed at the position of the other paired base. For example, if the i th base is paired with the j th base ($i < j$), then the i th position is represented by a left parenthesis and the j th position is represented by a right parenthesis, and then the element in the i th row and j th column as well as the element in the j th row and i th column of the base-base relationship matrix is set to 1. Therefore, for an RNA sequence with a length of 51, a secondary structure feature matrix X_{str} of 51×51 can be obtained.

2.2.3 Feature learning

In order to learn the low-dimensional representation of Nm sites, we take the extracted sequence features and chemical property features as the inputs of the convolutional neural network (CNN) (Wang et al., 2024b) respectively. This model is composed of an input layer, a convolutional layer, a pooling layer, and a fully connected layer. $X_{seq}^{(1)}$ and $X_{chem}^{(1)}$ respectively represent the convolutional features obtained from the convolutional layer (such as Equation 3):

$$\begin{aligned} X_{seq}^{(1)} &= \text{ReLU}(W_{seq} \otimes X_{seq} + b_{seq}) \\ X_{chem}^{(1)} &= \text{ReLU}(W_{chem} \otimes X_{chem} + b_{chem}) \end{aligned} \quad (3)$$

among them, X_{seq} and X_{chem} respectively represent the extracted sequence and physical and chemical property features. W_{seq} and W_{chem} represent the weight matrices of the convolutional kernels, b_{seq} and b_{chem} are the bias terms. $\otimes$ represents the convolution operation.

Then, the output of the convolutional layer goes through the pooling layer and the fully connected layer, and the finally obtained feature representation is as follows Equation 4:

$$\begin{aligned} X_{seq}^{(2)} &= f_{fc}((f_{pool}(f_{CNN}(X_{seq})))) \\ X_{chem}^{(2)} &= f_{fc}((f_{pool}(f_{CNN}(X_{chem})))) \end{aligned} \quad (4)$$

where $X_{seq}^{(2)}$ represents the sequence feature representation extracted by the convolutional neural network module, and $X_{chem}^{(2)}$ represents the chemical property feature representation extracted by the convolutional neural network module.

For the secondary structure features, we use the Graph Attention Network (GAT) (Veličković et al., 2018) for feature extraction. The GAT is composed of multiple stacked graph attention layers. Each graph attention layer contains multiple attention heads. Each attention head independently calculates the feature representation of nodes, and then combines these representations (such as concatenation or averaging) to obtain richer node features. Compared with traditional graph neural networks, the GAT has higher flexibility and expressive power.

The input of the graph attention layer is $h = \{h_1, h_2, \dots, h_N\}$, $h_i \in \mathbb{R}^F$, where N is the number of nodes and F is the number of features for each node. This layer generates a new set of node features $h' = \{h'_1, h'_2, \dots, h'_N\}$, $h'_i \in \mathbb{R}^{F'}$ as the output.

To obtain sufficient expressive power to transform the input features into higher-level features, at least one learnable linear transformation is required. For this purpose, as an initial step, a shared linear transformation with parameters $W \in \mathbb{R}^{F' \times F}$ is applied to each node. Then, self-attention of the nodes is performed, that is, a shared attention mechanism.

$$e_{ij} = \text{atten}(Wh_i, Wh_j) \quad (5)$$

among them, e_{ij} represents the importance of the feature of node j to node i . The attention mechanism $\text{atten}(\cdot)$ is a single-layer feedforward neural network, with the parameter being $\tilde{a} \in \mathbb{R}^{2F'}$.

In order to make the attention weights easy to compare among different nodes, we use the softmax function to standardize the selections for all j :

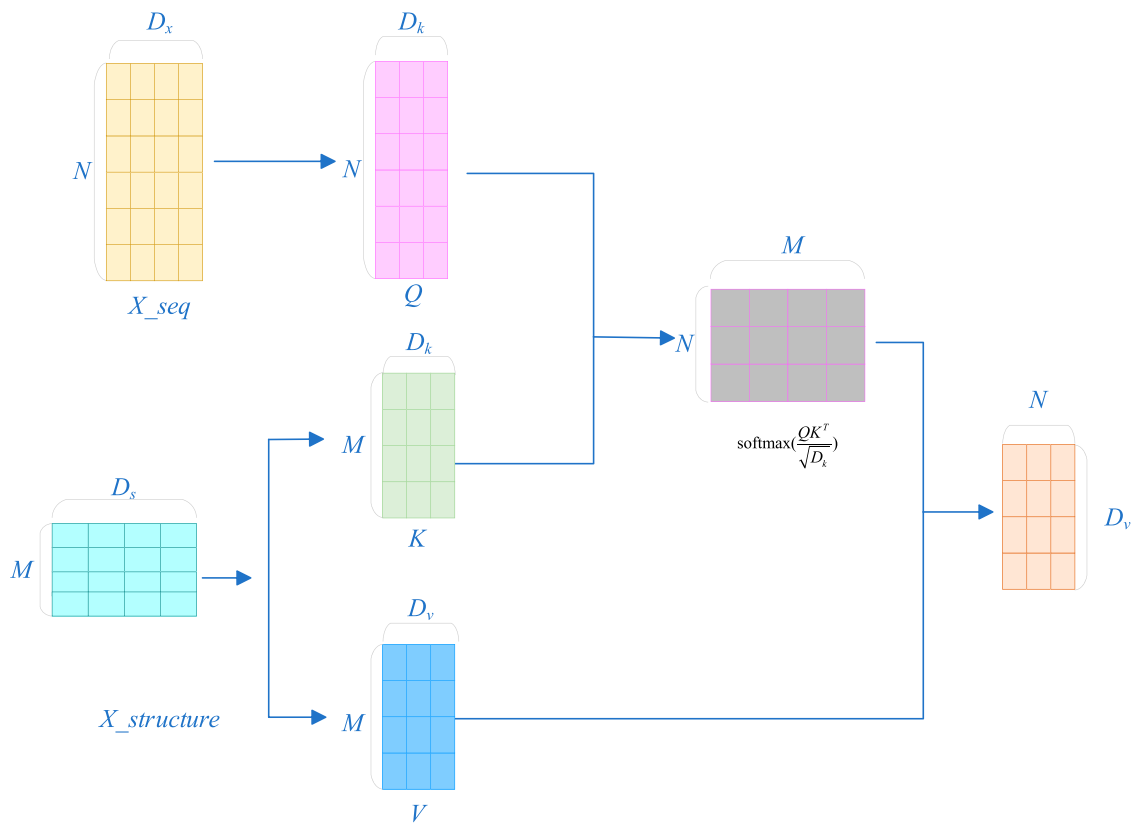


FIGURE 2
Flowchart of the cross-attention mechanism.

$$\alpha_{ij} = \text{softmax}_j(e_{ij}) = \frac{\exp(e_{ij})}{\sum_{k \in N_i} \exp(e_{ik})} \quad (6)$$

where N_i represents the neighborhood of node i .

Combining the above Equations 5, 6, the complete form of the attention mechanism can be written as:

$$\alpha_{ij} = \frac{\exp(\text{LeakyReLU}(\bar{a}^T [Wh_i \| Wh_j]))}{\sum_{k \in N_i} \exp(\text{LeakyReLU}(\bar{a}^T [Wh_i \| Wh_k]))} \quad (7)$$

among them, $\|$ represents the concatenation operation. Next combining Equation 7, the neighborhood representations of the nodes are linearly accumulated according to the attention weights to obtain the final output representation:

$$h'_i = \sigma \left(\sum_{j \in N_i} \alpha_{ij} Wh_j \right) \quad (8)$$

To stabilize the learning process and enrich the feature representation, GAT usually adopts the multi-head attention mechanism. By using K independent attention heads for calculation and then averaging the outputs of these heads, by modifying Equation 8, the final node features are obtained:

$$h'_i = \sigma \left(\frac{1}{K} \sum_{k=1}^K \sum_{j \in N_i} \alpha_{ij}^k W^k h_j \right) \quad (9)$$

where K is the number of attention mechanisms and W_k is the weight matrix for the k th attention mechanism.

Through the Graph Attention Network, we can obtain the deep representation $X_{str}^{(2)}$ of the secondary structure features of RNA through Equation 9.

2.2.4 Cross-attention module

The cross-attention module is a special attention mechanism used to process two different types of features. In this study, these two features are the sequence feature $X_{seq}^{(2)}$ and the secondary structure feature $X_{str}^{(2)}$ respectively. The cross-attention mechanism allows the model to adaptively focus on the relevant information in the other feature when processing one feature, thereby achieving effective fusion of the two features. The specific process is shown in Figure 2.

First, linear transformations are performed on the deep representation of sequence features $X_{seq}^{(2)}$ and the deep representation of secondary structure features $X_{str}^{(2)}$ respectively to obtain the query vector, key vector, and value vector respectively.

$$\begin{aligned} Q &= X_{seq}^{(2)} W_q \\ K &= X_{str}^{(2)} W_k \\ V &= X_{str}^{(2)} W_v \end{aligned} \quad (10)$$

where $W_q \in \mathbb{R}^{D_x \times D_k}$, $W_k \in \mathbb{R}^{D_y \times D_k}$, and $W_v \in \mathbb{R}^{D_y \times D_v}$ are learnable parameter matrices. $Q \in \mathbb{R}^{N \times D_k}$ is the query vector matrix,

$K \in \mathbb{R}^{M \times D_k}$ is the key vector matrix, and $V \in \mathbb{R}^{M \times D_v}$ is the value vector matrix.

Then, by calculating the similarity between the query vector and the key vector, the attention scores are obtained, and the scores are normalized to the interval $[0, 1]$ through the softmax function:

$$A = \text{softmax}\left(\frac{QK^T}{\sqrt{D_k}}\right) \quad (11)$$

where $A \in \mathbb{R}^{N \times M}$ is the attention score matrix.

Then, combining Equations 10, 11, the value vectors are weighted and summed according to the attention scores to obtain the fused features.

$$O = AV \quad (12)$$

where $O \in \mathbb{R}^{N \times D_v}$ is the fused feature matrix.

We concatenate $X_{chem}^{(2)}$ with the feature matrix O fused by the cross - attention module, and then input the result into two fully - connected layers to obtain the final prediction result.

2.2.5 Model training

We use binary cross entropy loss as the loss function, see Equation 13:

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (13)$$

where i represents the i th Nm, and y_i represents the true label, $\hat{y}_i$ represents the probability that the model predicts the class as positive. To minimize the loss function, we use the Adam optimizer (Kingma and Ba, 2014) to minimize the loss function.

2.3 Evaluation metrics

To evaluate the performance of the model, 5-fold cross-validation is used to evaluate the performance of the model. We plotted the Receiver Operating Characteristic (ROC) curve and the Precision-Recall curve, and calculated the Area Under the Curve of the ROC (AUC) and the Area Under the Precision-Recall Curve (AUPR) to assess the model's performance. The ROC curve is obtained by means of the True Positive Rate (TPR) and the False Positive Rate (FPR) at different scoring thresholds, and the Precision-Recall curve is obtained based on precision and recall at different scoring thresholds. The AUC is insensitive to whether the sample classes are balanced. In the case of highly imbalanced data, the performance is still overly ideal and cannot well reflect the actual situation. Under extremely imbalanced data (with fewer positive samples), the Precision-Recall (PR) curve may be more practical than the ROC curve. We used the AUC and AUPR as the main evaluation metrics. In addition, we adopted Accuracy (ACC), Matthews Correlation Coefficient (MCC), and $F1_score$ to present the results of the model, which are defined as follows Equation 14:

$$\begin{aligned} TPR &= \frac{TP}{TP + FN} \\ FPR &= \frac{FP}{FP + TN} \\ Precision &= \frac{TP}{TP + FP} \\ Recall &= \frac{TP}{TP + FN} \\ ACC &= \frac{TP + TN}{TN + FP + TP + FN} \\ MCC &= \frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TN + FN)(TP + FN)(TN + FP)}} \\ F1_score &= \frac{2 \times Precision \times Recall}{Precision + Recall} \end{aligned} \quad (14)$$

where TP represents the number of positive samples that are predicted as positive samples; FP represents the number of negative samples that are predicted as positive samples; TN represents the number of positive samples that are predicted as negative samples; FN represents the number of negative samples that are predicted as negative samples.

3 Results and discussion

3.1 Adjustment of parameters

In the MultiV_Nm model, we made fixed settings for some key parameters. Specifically, when using the convolutional neural network to process different features, the number of channels is determined according to the dimension of the features. That is, when processing sequence features, the number of channels is set to 4; when processing chemical property features, the number of channels is set to 3. Meanwhile, the size of the convolutional kernel is uniformly set to 3. For the graph attention network, we set the number of attention heads to 8 and the size of the max - pooling to 2. Next, we investigated one by one the impacts of the number of convolutional kernels in the convolutional neural network, the embedding dimension of the graph attention network (for the sake of consistency, we set the embedding dimension of the convolutional neural network to be the same as that of the graph attention network), the dropout rate, and the learning rate on the model's performance.

Inspired by TransRM (Liu et al., 2025), we set the number of convolutional kernels to 4, 8, 16, 32, and 64. As can be clearly observed from Figure 3A, as the number of convolutional kernels increases, the performance of the model shows a steady upward trend. It shows that the increase in the number of convolutional kernels enables the network to capture richer and more complex features. Based on this result, in this study, we defaulted the number of convolutional kernels to 64 to fully unleash the performance potential of the model.

Different embedding dimensions in the graph attention network have different impacts on the model's performance. In order to capture the key features of the graph data while avoiding overfitting, refer to GIAE-DTI (Wang et al., 2024a), we set the dimensions to gradually increase from 32 to 512, and attempt to find the

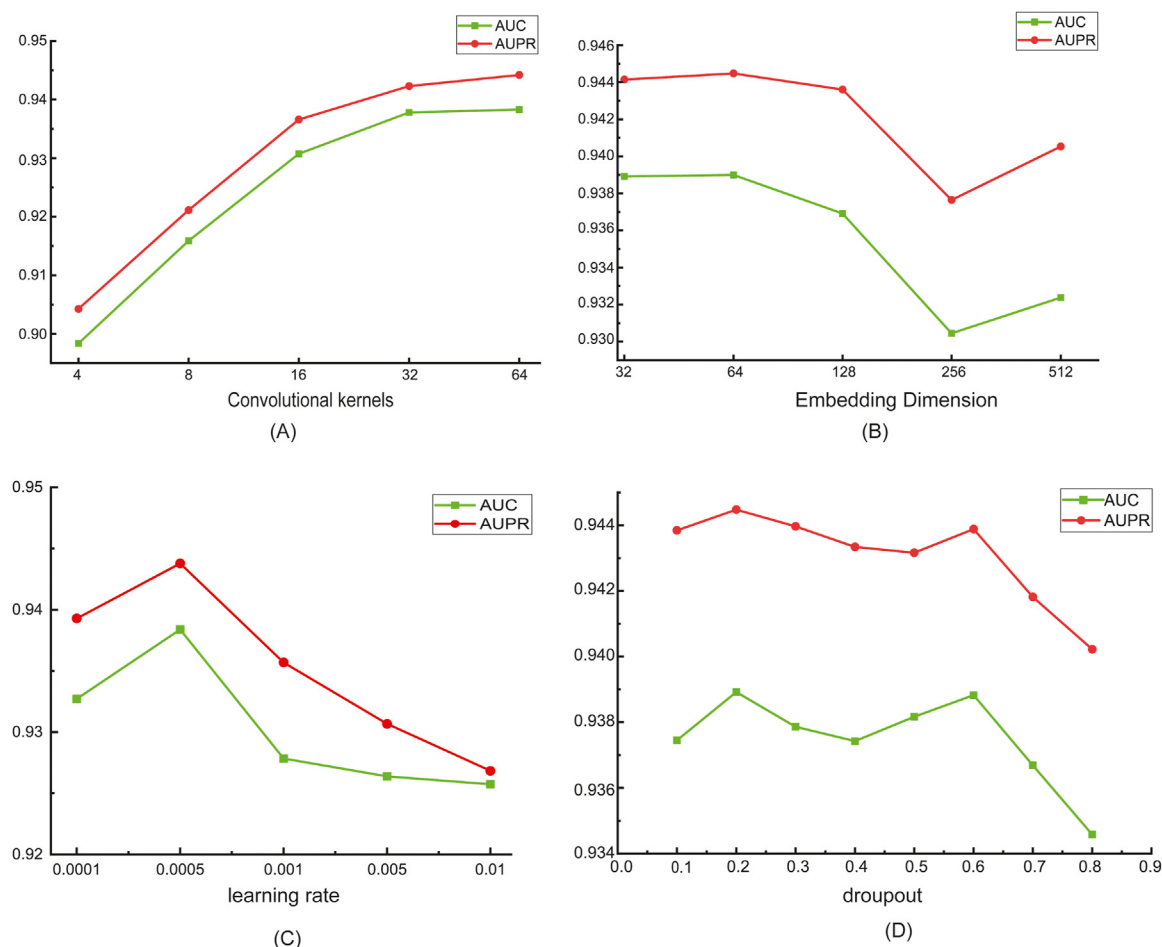


FIGURE 3
Parameter Sensitivity Analysis. (A) The number of convolutional kernels; (B) Embedding dimension; (C) Learning rate; (D) Dropout rate.

appropriate dimension that can fully describe the data features. We set the embedding dimensions as 32, 64, 128, 256 and 512. As shown in Figure 3B, as the embedding dimension increases, the AUC and AUPR values of the model first rise and then fall. When the embedding dimension is set to 64, the model achieves optimal performance. Since the secondary structure of the Nm-modified sequence is relatively simple, when the embedding dimension is too low, the model struggles to capture sufficient key information from the data. Conversely, when the embedding dimension is too high, it leads to excessive information redundancy, increasing the computational burden and potentially introducing noise interference. Therefore, we set the default value of the embedding dimension to 64 to balance model performance and information utilization efficiency.

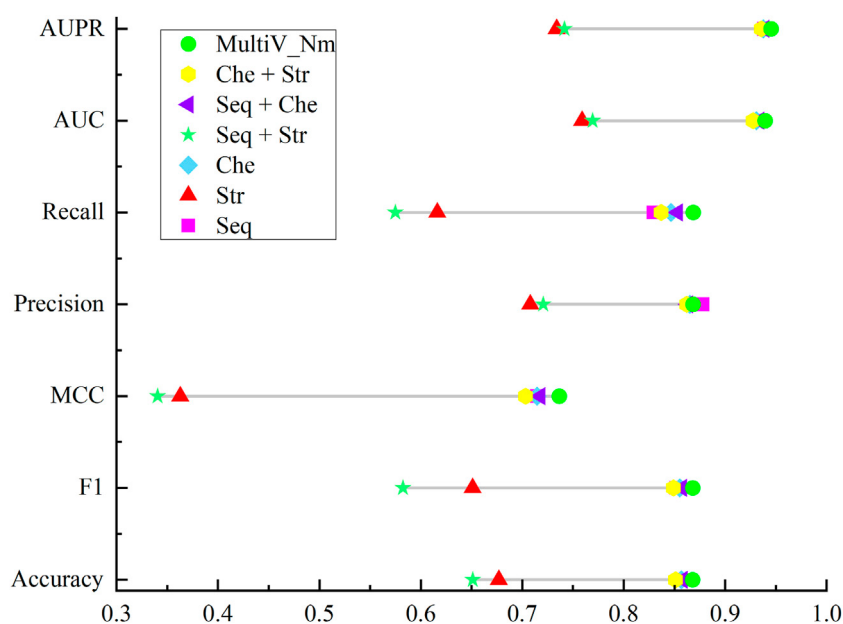
It is evident from Figure 3C that the learning rate has a significant impact on the experimental results. Refer to DualC (Guo et al., 2025), we experimented with several different learning rate values, including 0.0001, 0.0005, 0.001, 0.005, and 0.01. The experimental results show that when the learning rate is set to 0.0005, the model exhibits optimal performance. This fully demonstrates that the learning rate, as a hyperparameter, plays a crucial regulatory role in the model training process. When the learning rate is too small, the step size of parameter updates during

model training is too short, resulting in an extremely slow convergence rate. The model may require a large number of training epochs to achieve good performance, failing to fully learn the effective features in the data. On the other hand, when the learning rate is too large, the step size of parameter updates is too big, causing the model difficultly to converge, leading to poor performance. When the learning rate is set to 0.0005, the model has a moderate parameter update step size, thus achieving the best performance.

During the model training process, overfitting is a common problem, which causes the model to perform excellently on the training set but have poor generalization ability on the test set or new data. Dropout, as a simple and effective regularization technique, can significantly alleviate this problem. In this experiment, to investigate the impact of the dropout rate on model performance, refer to GIAE-DTI (Wang et al., 2024a), we set dropout rates as 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, and 0.8. Through in-depth analysis of the experimental data in Figure 3D, we found that the change in the dropout rate has a very significant impact on model performance. When the dropout rate is set too low, the model cannot effectively suppress the co-adaptation between neurons, and the overfitting problem remains severe. When the dropout rate is too high, the model randomly discards too many neurons, resulting in a large loss

TABLE 2 Feature combinations in the ablation study.

Method name	Sequence	Chemical property	Second structure
Seq	✓		
Che		✓	
Str			✓
Seq + Che	✓	✓	
Seq + Str	✓		✓
Che + Str		✓	✓
MultiV_Nm	✓	✓	✓

FIGURE 4
Comparison of the results of the ablation experiment.

of useful information learned by the model and causing underfitting. When the dropout rate is set to 0.2, the model can effectively prevent overfitting while retaining enough useful information, fully learn the feature patterns of the data, and significantly improve the model's generalization ability. Based on this, in the subsequent model training and optimization process of this study, we fixed the dropout rate at 0.2.

3.2 Ablation study

When predicting Nm methylation modification sites, the MultiV_Nm model integrates sequence features, chemical property features, and secondary structure features. These features reflect the characteristics of biomolecules from different dimensions. The sequence features contain the genetic information of base arrangement, the chemical property features reflect the chemical properties of molecules, and the secondary structure features describe the spatial folding

morphology of molecules. Together, they provide support for accurate prediction.

To deeply analyze the specific roles of these three features in the model, a feature ablation experiment was constructed. By removing one or two features respectively, we compared the performance of the simplified model with that of the original model that fuses the three features. Table 2 clearly lists the feature combinations used in each experimental method.

The results of the ablation experiment are shown in Figure 4. As can be seen from Figure 4, when the three features are used separately, the prediction effects of only the sequence features and the chemical property features are not very different, while the prediction effect of using the secondary structure features is the worst. This indicates that in the scenario of predicting Nm methylation sites, sequence information and chemical properties can more effectively characterize the features of Nm methylation sites. In contrast, the secondary structure has obvious limitations. On the one hand, the secondary structure analysis focuses on the spatial conformation of RNA and cannot reflect the specific

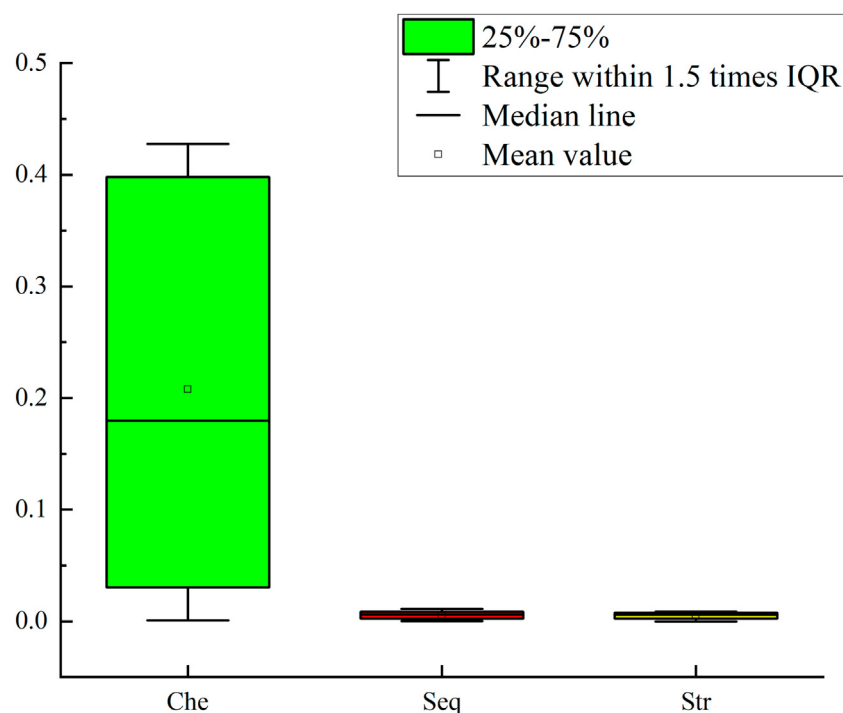


FIGURE 5
Boxplots of the importance of three types of features.

composition of the sequence, making it difficult to capture the local features in the sequence. On the other hand, when the sequence to be analyzed is too short, there will be large errors in the predicted secondary structure, resulting in the inability to provide reliable support for the prediction of Nm methylation sites. In addition, different combinations of features have different effects on the prediction of Nm methylation sites. Through research, it has been found that in the case of pairwise feature combinations, the combination of sequence features and chemical property features has the most prominent prediction effect; the combination of chemical property features and secondary structure features has a slightly inferior prediction effect. It is worth noting that the Multi_Nm model, by fusing the three features of sequence features, chemical property features, and secondary structure, has demonstrated the most excellent performance in predicting methylation sites, further improving the accuracy and reliability of the prediction.

To further illustrate the importance of each feature for the model's prediction, we evaluated the importance of the three types of features using a permutation-based feature importance calculation method. Taking the sequence features as an example, each time the order of one sequence feature was shuffled. Through five-fold cross-validation, we calculated the difference between the *AUC* of the model based on the permuted feature and the *AUC* of the standard model to obtain the importance of the sequence feature. The processing method for chemical property features and secondary structure features was the same as that for sequence features, and an importance vector with 51 dimensions was obtained for each type of feature. We plotted the importance of the three types of features using boxplots. As can be seen from Figure 5, when the

chemical property features were permuted, the model was affected to the greatest extent, while the impacts of the sequence features and secondary structure features were relatively small. However, the degree of influence of the sequence features on the model was still higher than that of the secondary structure features.

3.3 MultiV_Nm performance under 5-CV and 10-CV

Cross-validation is an important means to evaluate the generalization ability of a model. Here, we demonstrate the performance of the MultiV_Nm model in the scenarios of five-fold cross-validation (5CV) and ten-fold cross-validation (10CV). In five-fold cross-validation, the dataset is evenly divided into five parts. One part is taken as the test set in turn, and the remaining four parts are used as the training set. The training and testing processes are repeated five times. Similarly, in ten-fold cross-validation, the dataset is divided into ten parts for operation.

Figure 6 plots the ROC curve and the PR curve generated during the 5CV process. Figure 7 presents the performance of the model under 10CV. Through analysis, it is found that, whether in the case of 5CV or 10CV, the fluctuation range of the curves obtained from each fold of validation is extremely small. At the same time, the *AUC* and *AUPR* of the model both remain at a relatively high level. This fully demonstrates that the MultiV_Nm model has strong robustness and stability, and is capable of maintaining good generalization ability and prediction accuracy under different data distributions.

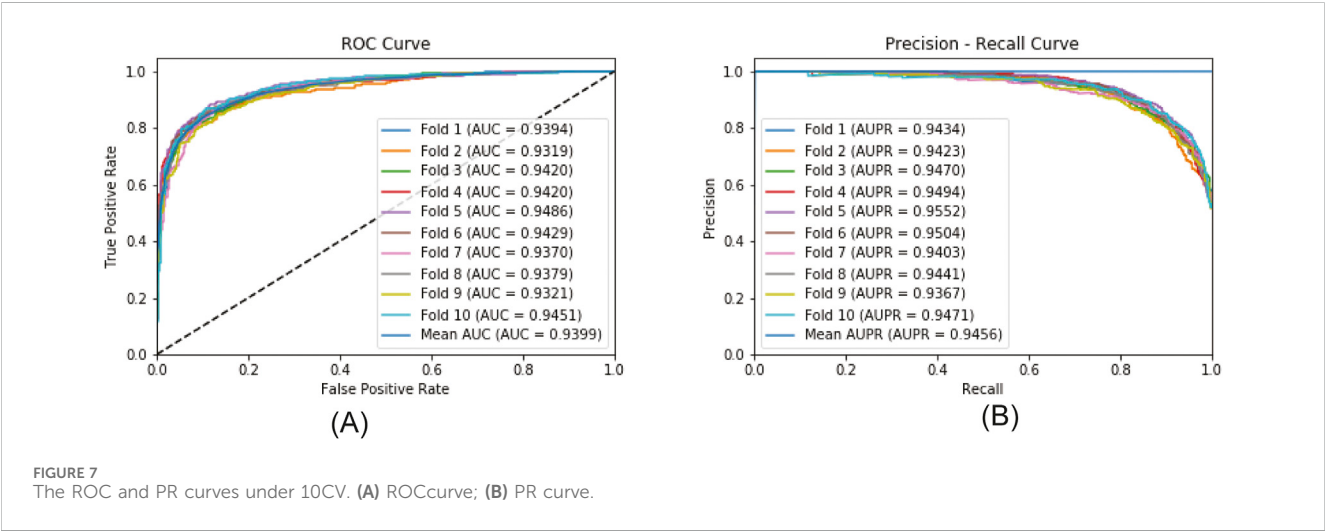
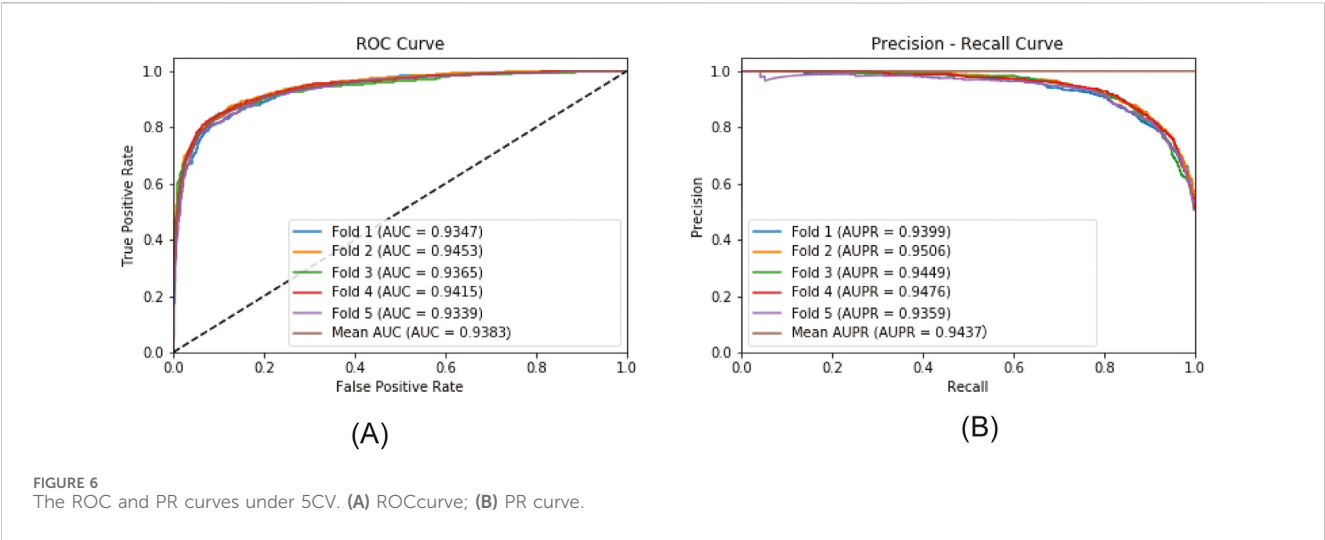
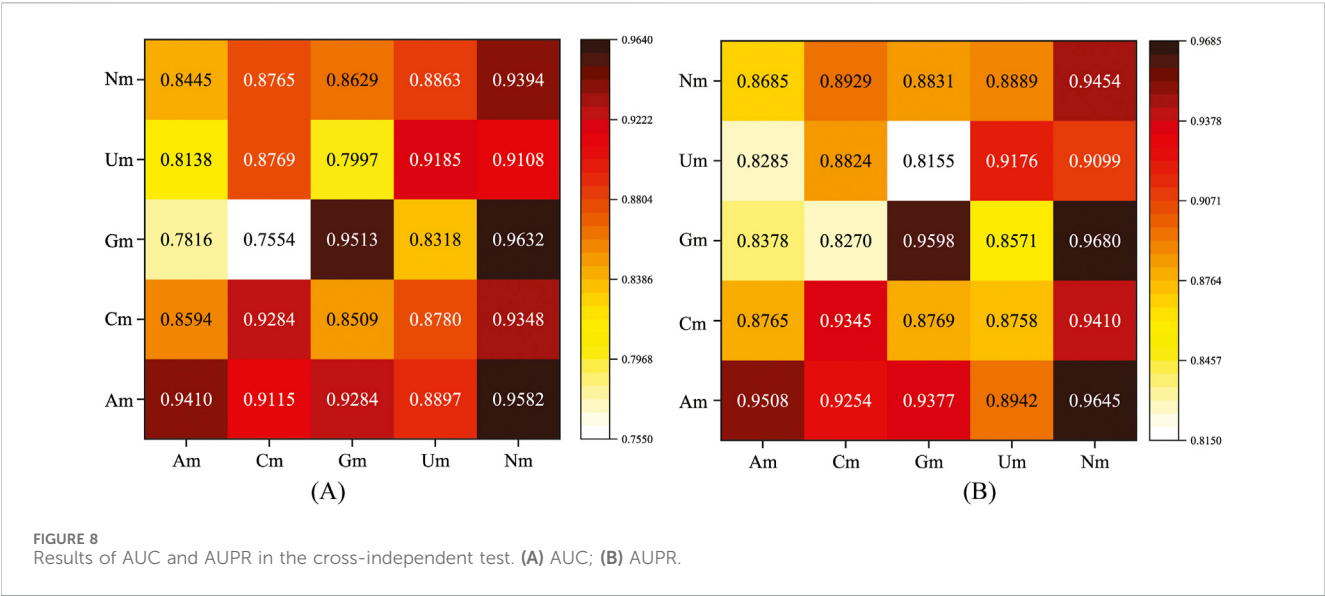


TABLE 3 Comparison of the results of multiple experiments.

Cross-validation	ACC	F1_score	Precision	Recall	MCC	AUC	AUPR
5CV (single round)	0.8587	0.8582	0.8629	0.8563	0.7198	0.9383	0.9437
5CV (10 rounds)	0.8639	0.8630	0.8701	0.8583	0.7297	0.9389	0.9445
	±0.0072	±0.0079	±0.0218	±0.0284	±0.0135	±0.0045	±0.0044
10CV (single round)	0.8647	0.8641	0.8700	0.8604	0.7312	0.9399	0.9456
10CV (10 rounds)	0.8678	0.8668	0.8747	0.8608	0.7372	0.9405	0.9463
	±0.0116	±0.0116	±0.0260	±0.0289	±0.0224	±0.0071	±0.0066

To further evaluate the robustness of the model, statistical methods were used for analysis. Ten rounds of 5CV and 10CV were respectively carried out, and their means and standard deviations were calculated for analysis. The specific results are shown in Table 3. As can be seen from Table 3, the results of the ten experiments are quite close to a single experiment. In the ten

experiments, regardless of whether it is 5CV or 10CV, the standard deviation of each indicator is less than 0.05, while the standard deviations of AUC and AUPR are both less than 0.01. This indicates that MultiV_Nm has excellent robustness and is minimally affected by the randomness of the divided dataset.



3.4 Cross-independent testing

Nm methylation modifications mainly include four types, namely, Am, Cm, Gm, and Um. Next, we conducted cross-independent tests using the MultiV_Nm model. Specifically, the model was trained using four single types of Nm (i.e., Am, Cm, Gm, and Um) and the total Nm containing all types. Then, independent test sets were used respectively to evaluate the performance of the trained model. The evaluation metrics selected were the *AUC* and *AUPR*. The experimental results are shown in Figure 8.

In Figure 8, the horizontal axis represents the data types used for model training, and the vertical axis represents the data types used for testing. As can be clearly seen from the data presented in the chart, when the total Nm is used for model training, regardless of which single type of Nm is used for testing, the model can achieve relatively ideal prediction results. When a single type of Nm is used for training, the model can only achieve the optimal prediction performance when the corresponding type of Nm is used for testing.

Through in-depth analysis of these experimental results, we can infer that the total Nm contains the information of all types of Nm. This enables the model to fully learn the characteristic patterns shared by different types of Nm during the training process, thus possessing a broader adaptability and generalization ability. Therefore, when testing a single type of Nm, it can demonstrate good prediction performance. When a single type of Nm is used for training, the characteristic patterns learned by the model are highly matched to that specific type of Nm. So, when predicting the same type of Nm, due to the consistency of the characteristic patterns, the model can more accurately capture the patterns in the data, thereby obtaining better prediction results.

3.5 Comparison with existing methods

To comprehensively verify the effectiveness of the algorithm adopted by the MultiV_Nm model, we carried out comparative tests. The MultiV_Nm model was compared with NmRF based on

TABLE 4 Performance comparison for Nm methylation sites prediction.

Method	Accuracy	F1	MCC	Precision	Recall
NmRF	0.5419	0.6299	0.0953	0.5284	0.7797
Deep-2'-O-Me	0.6580	0.7046	0.3331	0.6201	0.8160
MultiV-Nm	0.8679	0.8681	0.7365	0.8683	0.8686

machine learning and Deep-2'-O-Me based on deep learning in independent test. NmRF relies on the website <http://lab.malab.cn/~acy/NmRF> to provide prediction services. When predicting Nm methylation sites, this website only outputs the prediction results and cannot provide more derivative data. To ensure the consistency and fairness of the comparative tests, we selected *Precision*, *Recall*, *ACC*, *MCC*, and *F1_score* to quantitatively evaluate the prediction performance of each model. In addition, when using Deep-2'-O-Me to test, the predicted probabilities of all sites are less than 0.5. After repeated experiments and analysis, in order to enable this method to effectively output results, we set the threshold to 0.3.

As can be seen from Table 4, the MultiV_Nm model significantly outperforms the NmRF and Deep-2'-O-Me methods in various evaluation indicators. The *ACC* of the MultiV_Nm reaches 0.8679, which has a very prominent advantage compared with 0.5419 of the NmRF and 0.6580 of the Deep-2'-O-Me. In terms of the *MCC* indicator, 0.7365 of the MultiV_Nm is much higher than 0.0953 of the NmRF and 0.3331 of the Deep-2'-O-Me. The experimental results show that the MultiV_Nm has excellent performance in the prediction of Nm methylation sites.

4 Conclusion

In the realm of RNA modification research, Nm methylation modification is pivotal. It participates in key biological processes. Precise identification of Nm methylation sites aids in uncovering disease mechanisms and developing novel diagnostic and treatment

strategies. In this paper, we proposed MultiV_Nm, a multi-view feature - based prediction framework for 2'-O methylation sites. On the basis of separately extracting the sequence features, chemical features, and secondary structure features of Nm methylation sites, we used a convolutional neural network and a graph attention network, and combined them with a cross-attention mechanism to predict the Nm methylation sites. Compared with existing methods, MultiV_Nm performs excellently in multiple evaluation indicators.

However, MultiV_Nm still has some limitations. First, the model relies on high-quality RNA modification data. When extracting secondary structure features, if the data accuracy is low, it will lead to inaccurate secondary structure prediction, thereby reducing the prediction accuracy. Second, this study only used information on human Nm modification sites and did not extend it to the prediction of Nm modification sites in other species. Third, although MultiV_Nm can be extended to the prediction of other types of RNA modification sites, for different types of RNA modifications, it may be necessary to redesign the feature extraction methods and model structure to adapt to their unique modification patterns and characteristics. Even so, MultiV_Nm can still provide new ideas and insights for the prediction of RNA modification sites in different species and different types, helping to promote basic research and potentially bringing breakthroughs in the field of biomedicine.

Data availability statement

The original contributions presented in the study are included in the article/[Supplementary Material](#), further inquiries can be directed to the corresponding authors.

Author contributions

LB: Data curation, Writing – review and editing, Software, Methodology, Conceptualization, Validation, Writing – original draft. FL: Conceptualization, Supervision, Writing – review and editing, Funding acquisition. YW: Writing – review and editing, Data curation, Software, Methodology. JS: Data curation, Methodology, Writing – review and editing, Software. LL: Data

curation, Writing – review and editing, Funding acquisition, Supervision, Conceptualization, Investigation.

Funding

The author(s) declare that financial support was received for the research and/or publication of this article. This work was supported by the Natural Science Basic Research Program of Shaanxi (2024JC-YBQN-0624); The Fundamental Research Funds for the Central Universities, Shaanxi Normal University (GK202406008); The National Natural Science Foundation of China (62402010).

Conflict of interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Generative AI statement

The author(s) declare that no Generative AI was used in the creation of this manuscript.

Publisher's note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

Supplementary material

The Supplementary Material for this article can be found online at: <https://www.frontiersin.org/articles/10.3389/fgene.2025.1608490/full#supplementary-material>

References

- Agris, P. F., Narendran, A., Sarachan, K., Yare, V. Y. P., and Eruysal, E. (2017). The Importance of being modified: the role of RNA modifications in translational fidelity. *Enzymes* 41, 1–50. doi:10.1016/bs.enz.2017.03.005
- Ao, C., Zou, Q., and Yu, L. (2022). NmRF: identification of multispecies RNA 2'-O-methylation modification sites from RNA sequences. *Briefings Bioinforma.* 23 (1), bbab480. doi:10.1093/bib/bbab480
- Blijlevens, M., Li, J., and van Beusechem, V. W. (2021). Biology of the mRNA splicing machinery and its dysregulation in cancer providing therapeutic opportunities. *Int. J. Mol. Sci.* 22, 5110. doi:10.3390/ijms22105110
- Boccalletto, P., Machnicka, M. A., Purta, E., Pawel, P., Bagiński, B., Wirecki, T. K., et al. (2018). MODOMICS: a database of RNA modification pathways. 2017 update. *Nucleic Acids Res.* 46 (D1), D303–D307. doi:10.1093/nar/gkx1030
- Boccalletto, P., Stefaniak, F., Ray, A., Cappannini, A., Mukherjee, S., Purta, E., et al. (2022). MODOMICS: a database of RNA modification pathways. 2021 update. *Nucleic Acids Res.* 50 (D1), D231–D235. doi:10.1093/nar/gkab1083
- Chen, W., Feng, P., Tang, H., Ding, H., and Lin, H. (2016). Identifying 2'-O-methylation sites by integrating nucleotide chemical properties and nucleotide compositions. *Genomics* 107 (6), 255–258. doi:10.1016/j.ygeno.2016.05.003
- Cui, L., Ma, R., Cai, J., Guo, C., Chen, Z., Yao, L., et al. (2022). RNA modifications: importance in immune cell biology and related diseases. *Signal Transduct. Target. Ther.* 7, 334. doi:10.1038/s41392-022-01175-9
- Daffis, S., Szretter, K. J., Schriewer, J., Li, J., Diamond, M. S., Errett, J., et al. (2010). 2'-O methylation of the viral mRNA cap evades host restriction by IFIT family members. *Nature* 468 (7322), 452–456. doi:10.1038/nature09489
- Dai, Q., Moshitch-Moshkovitz, S., Han, D., Kol, N., Amariglio, N., Rechavi, G., et al. (2017). Nm-seq maps 2'-O-methylation sites in human mRNA with base precision. *Nat. Methods* 14 (7), 695–698. doi:10.1038/nmeth.4294
- Guo, L., Lei, X., Liu, L., Chen, M., and Pan, Y. (2025). DualC: drug-drug interaction prediction based on dual latent feature extractions. *IEEE Trans. Emerg. Top. Comput. Intell.* 9 (1), 946–960. doi:10.1109/tetci.2024.3502414

- Kingma, D. P., and Ba, J. (2014). "Adam: a method for stochastic optimization," in International conference on learning representations.
- Li, H., Chen, L., Huang, Z., Luo, X., Li, H., Ren, J., et al. (2021). DeepOME: a web server for the prediction of 2'-O-Me sites based on the hybrid CNN and BLSTM architecture. *Front. Cell Dev. Biol.* 9, 686894. doi:10.3389/fcell.2021.686894
- Li, Y. H., Yu, C. Y., Li, X. X., Zhang, P., Tang, J., Yang, Q., et al. (2018). Therapeutic target database update 2018: enriched resource for facilitating bench-to-clinic research of targeted therapeutics. *Nucleic Acids Res.* 46 (D1), D1121–D1127. doi:10.1093/nar/gkx1076
- Li, Y. Q., Yi, Y., Gao, X. L., Zhang, L. L., Cao, Q., Chen, K. F., et al. (2024). 2'-O-methylation at internal sites on mRNA promotes mRNA stability. *Mol. Cell* 84 (12), 2320–2336.e6. doi:10.1016/j.molcel.2024.04.011
- Lim, S. L., Qu, Z. P., Kortschak, R. D., Lawrence, D. M., Geoghegan, J., Hempfling, A. L., et al. (2015). Correction: HENMT1 and piRNA stability are required for adult male germ cell transposon repression and to define the spermatogenic program in the mouse. *PLoS Genet.* 11, e1005782. doi:10.1371/journal.pgen.1005782
- Liu, L., Lei, X., Fang, Z., Tang, Y., Meng, J., and Wei, Z. (2020). ISGM1A: integration of sequence features and genomic features to improve the prediction of human m1A RNA methylation sites. *IEEE Access* 8, 81971–81977. doi:10.1109/access.2020.2991070
- Liu, L., Lei, X., Wang, Z., Meng, J., and Song, B. (2025). TransRM: weakly supervised learning of translation-enhancing N6-methyladenosine (m6A) in circular RNAs. *Int. J. Biol. Macromol.* 306, 141588. doi:10.1016/j.ijbiomac.2025.141588
- Lorenz, R., Bernhart, S. H., Siederdisen, C. H. Z., Tafer, H., Flamm, C., Stadler, P. F., et al. (2011). ViennaRNA package 2.0. *Algorithms Mol. Biol.* 6 (1), 26. doi:10.1186/1748-7188-6-26
- Maden, B. E., Corbett, M. E., Heeney, P. A., Pugh, K., and Ajuh, P. M. (1995). Classical and novel approaches to the detection and localization of the numerous modified nucleotides in eukaryotic ribosomal RNA. *Biochimie* 77 (1-2), 22–29. doi:10.1016/0300-9084(96)88100-4
- Marcel, V., Ghayad, S. E., Belin, S., Therizols, G., Morel, A. P., Gonzalez, E. S., et al. (2013). p53 acts as a safeguard of translational control by regulating fibrillarin and rRNA methylation in cancer. *Cancer Cell* 24, 318–330. doi:10.1016/j.ccr.2013.08.013
- Mostavi, M., Salekin, S., and Huang, Y. (2018). Deep-2'-O-me: predicting 2'-O-methylation sites by convolutional neural networks. *Annu. Int. Conf. IEEE Eng. Med. Biol. Soc.*, 2394–2397. doi:10.1109/EMBC.2018.8512780
- Nicolai, K., Jansson, M. D., Häfner, S. J., Disa, T., Ulf, B., Mikkil, C. D., et al. (2016). Profiling of 2'-O-Me in human rRNA reveals a subset of fractionally modified positions and provides evidence for ribosome heterogeneity. *Nucleic Acids Res.* 44 (16), 7884–7895. doi:10.1093/nar/gkw482
- Picard-Jean, F., Brand, C., Tremblay-Letourneau, M., Allaire, A., Beaudoin, M. C., Boudreault, S., et al. (2018). Correction: 2'-O-methylation of the mRNA cap protects RNAs from decapping and degradation by DXO. *PLoS One* 13, e0202308. doi:10.1371/journal.pone.0202308
- Pichot, F., Marchand, V., Helm, M., and Motorin, Y. (2022). Machine learning algorithm for precise prediction of 2'-O-methylation (Nm) sites from experimental RiboMethSeq datasets. *Methods* 203, 311–321. doi:10.1016/j.ymeth.2022.03.007
- Qiu, W. R., Jiang, S. Y., Sun, B. Q., Xiao, X., Cheng, X., and Chou, K. C. (2017). iRNA-2methyl: identify RNA 2'-O-methylation sites by incorporating sequence-coupled effects into general PseKNC and ensemble classifier. *Med. Chem.* 13 (8), 734–743. doi:10.2174/1573406413666170623082245
- Ringard, M., Marchand, V., Decroly, E., Motorin, Y., and Bennasser, Y. (2019). FTSJ3 is an RNA 2'-O-methyltransferase recruited by HIV to avoid innate immune sensing. *Nature* 565, 500–504. doi:10.1038/s41586-018-0841-4
- Roundtree, I. A., Evans, M. E., Pan, T., and He, C. (2017). Dynamic RNA modifications in gene expression regulation. *Cell* 169 (7), 1187–1200. doi:10.1016/j.cell.2017.05.045
- Song, B., Wang, X., Liang, Z., Ma, J., Huang, D., Wang, Y., et al. (2023). RMDisease V2.0: an updated database of genetic variants that affect RNA modifications with disease and trait implication. *Nucleic Acids Res.* 51, D1388–D1396. doi:10.1093/nar/gkac750
- Soylu, N. N., and Sefer, E. (2023). BERT2OME: prediction of 2'-O-methylation modifications from RNA sequence by transformer architecture based on BERT. *IEEE/ACM Trans. Comput. Biol. Bioinforma.* 20, 2177–2189. doi:10.1109/TCBB.2023.323769
- Tahir, M., Tayara, H., and Chong, K. T. (2019). iRNA-PseKNC(2methyl): identify RNA 2'-O-methylation sites by convolution neural network and Chou's pseudo components. *J. Theor. Biol.* 465, 1–6. doi:10.1016/j.jtbi.2018.12.034
- Veličković, P., Cucurull, G., Casanova, A., Romero, A., Liò, P., and Bengio, Y. (2018). "Graph attention networks," in International conference on learning representations (ICLR).
- Wang, M., Lei, X., Liu, L., Chen, J., and Wu, F. (2024a). GIAE-DTI: predicting drug-target interactions based on heterogeneous network and GIN-based graph autoencoder based on heterogeneous network and GIN-based graph autoencoder. *IEEE J. Biomed. Health Inf.*, 1–14. doi:10.1109/JBHI.2024.3458794
- Wang, Z., Lei, X., and Pan, Y. (2024b). An encoding-decoding framework based on CNN for circRNA-RBP binding sites prediction. *Chin. J. Electron.* 33, 256–263. doi:10.23919/cje.2022.00.361
- Wu, K., Li, Y., Yi, Y., Yu, Y., Wang, Y., Zhang, L., et al. (2024). The detection, function, and therapeutic potential of RNA 2'-O-methylation. *Innovation Life* 3, 100112. doi:10.59717/j.xinn-life.2024.100112
- Xuan, J. J., Sun, W. J., Lin, P. H., Zhou, K. R., Liu, S., Zheng, L. L., et al. (2018). RMBase v2.0: deciphering the map of RNA modifications from epitranscriptome sequencing data. *Nucleic Acids Res.* 46 (D1), D327–D334. doi:10.1093/nar/gkx934
- Yang, H., Lv, H., Ding, H., Chen, W., and Lin, H. (2018). iRNA-2OM: a sequence-based predictor for identifying 2'-O-methylation sites in *Homo sapiens*. *J. Comput. Biol.* 25 (11), 1266–1277. doi:10.1089/cmb.2018.0004
- Yang, Y. H., Ma, C. Y., Gao, D., Liu, X. W., Yuan, S. S., and Ding, H. (2023). i2OM: toward a better prediction of 2'-O-methylation in human RNA. *Int. J. Biol. Macromol.* 239, 124247. doi:10.1016/j.ijbiomac.2023.124247
- Yu, Y. T., Shu, M. D., and Steitz, J. A. (1997). A new method for detecting sites of 2'-O-methylation in RNA molecules. *RNA* 3, 324–331.
- Yu, Y. T., Terns, R. M., and Terns, M. P. (2004). "Mechanisms and functions of RNA-guided RNA modification," in *Mechanisms and functions of RNA-guided RNA modification*, 12. Springer Berlin Heidelberg, 223–262. doi:10.1007/b105585
- Zhang, P., Huang, J., Zheng, W., Chen, L., Liu, S., Liu, A., et al. (2023). Single-base resolution mapping of 2'-O-methylation sites by an exoribonuclease-enriched chemical method. *Sci. China Life Sci.* 66 (4), 800–818. doi:10.1007/s11427-022-2210-0
- Zhou, Y., Cui, Q., and Zhou, Y. (2019). NmSEER V2.0: a prediction tool for 2'-O-methylation sites based on random forest and multi-encoding combination. *BMC Bioinforma.* 20 (1), 690. doi:10.1186/s12859-019-3265-8