

A dual diffusion model-based representation learning framework for antimicrobial peptides classification

Wen Kong¹, Lingling Fu¹, Xingpeng Jiang^{1,2,3,*}, Weizhong Zhao^{1,2,3,*} 

¹Hubei Provincial Key Laboratory of Artificial Intelligence and Smart Learning, Central China Normal University, Wuhan, Hubei 430079, China

²School of Computer, Central China Normal University, Wuhan, Hubei 430089, China

³National Language Resources Monitoring & Research Center for Network Media, Central China Normal University, Wuhan, Hubei 430089, China

*Corresponding authors. Xingpeng Jiang, School of Computer, Central China Normal University, Wuhan, Hubei 430089, China. E-mail: xpijiang@mail.ccnu.edu.cn; Weizhong Zhao, School of Computer, Central China Normal University, Wuhan, Hubei 430089, China. E-mail: wzzhao@ccnu.edu.cn.

Associate Editor: Macha Nikolski

Abstract

Motivation: The increasing prevalence of antibiotic-resistant bacteria has intensified the demand for novel antimicrobial agents. Antimicrobial peptides (AMPs) have emerged as promising alternatives, yet their identification or classification remains challenging due to the lack of multi-perspective information, insufficient feature representation learning, and monocular data modalities.

Results: In this paper, we propose a dual diffusion model-based representation learning framework for classifying AMPs, which effectively integrates both peptide sequence and structure information to address existing issues for the task. Specifically, our approach utilizes a multi-view feature construction module, which encodes peptide sequences and structures from distinctive perspectives, deriving initial feature representations with enriched biological semantics. To enhance representation learning, the proposed framework leverages both diffusion models for sequence and structure information respectively to effectively capture complex semantics from dual modalities. In addition, both single-modal and dual-modal contrastive learning are used to further advance the representation learning. Results of comprehensive experiments demonstrate that our model outperforms existing methods for the task of AMPs classification, providing a feasible solution to accelerating the discovery of novel antimicrobial agents.

Availability of implementation: The data and source codes are available in GitHub at <https://github.com/kww567upup/DDM>.

1 Introduction

Antimicrobial resistance (AMR) to antibiotics has rapidly evolved, posing major challenges to medicine and public health (Salam *et al.* 2023). Antimicrobial peptides (AMPs), naturally occurring short peptides, exhibit strong biofilm disruption, lower resistance development, and broad-spectrum antimicrobial activity, making AMPs prediction crucial for novel antimicrobial therapy development (Yan *et al.* 2022). However, complex sequence-structure-function relationships hinder rapid screening from large datasets, and the lack of sufficient biological information in modeling peptides complicates the therapeutic development.

Deep learning (DL), with its powerful pattern recognition and feature extraction capabilities, can more efficiently perform various prediction tasks from large datasets. Consequently, researchers have applied DL to AMPs recognition and classification, to accelerate the development of drugs to combat antibiotic resistance (Al-Omari *et al.* 2024). Existing AMPs recognition methods can be broadly categorized into three types: sequence-based

methods, structure-based methods, and hybrid methods that integrate both types of information. First, sequence-based methods typically analyze amino acid composition, physicochemical properties, *k*-mer features, or substitution matrices (e.g. BLOSUM, PAM, PC6) to derive similarity scores for AMPs prediction (Henikoff and Henikoff 1992, Mount 2008). In recent years, in addition to analyzing the biological characteristics of amino acid sequences for sequence modeling, researchers have started to use knowledge from the field of natural language processing (NLP) to design protein language models capable of uncovering latent patterns in AMPs sequences, such as ProtBERT (Brandes *et al.* 2022) and ESM series (Lin *et al.* 2022). However, existing sequence-based models primarily utilize sequence information from relatively singular perspectives, resulting in suboptimal quality of obtained representations. Conversely, structure-based methods focus on the critical determinant of functional activity, i.e. peptide structure, to derive more meaningful representations. These methods analyze spatial structure information using graph representation networks. For example, in SGAC, residue-level $C\alpha$

Received: 11 September 2025. Revised: 24 January 2026. Accepted: 10 February 2026

© The Author(s) 2026. 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.

distance graph is constructed for structural embedding, and related atom-bond graph is modeled to represent peptides through molecular connectivity (Wang *et al.* 2026). In AMPs-Net, sequences are converted into atom-bond graphs with physicochemical atom and bond attributes for classification (Ruiz Puentes *et al.* 2022). PepMNet introduces hierarchical aggregation from atoms to residues to better accommodate non-canonical chemical features (Garzon Otero *et al.* 2025). Nevertheless, these methods rely on the single structure-derived graph as a deterministic scaffold, which limits robustness to structure prediction errors and conformational variability. Thereafter, some hybrid methods are proposed by combining both sequence and structure information to improve representation learning, such as SAMPpred-GAT (Yan *et al.* 2023) and deepAMPNet (Zhao *et al.* 2024). These methods typically utilize predicted structural contacts to define a fixed graph adjacency matrix, and use sequence-derived features as node attributes for graph representation learning.

Generally speaking, despite significant progress in accurate classification of AMPs, two critical challenges remain unresolved. First, current approaches lack comprehensive integration of heterogeneous data from multiple perspectives. This limitation results in insufficient utilization of multi-source heterogeneous data, failing to fully exploit their complementary advantages (Alizadehsani *et al.* 2024). Second, feature extraction in representation learning modules remains inadequate. The key to improving AMPs identification accuracy lies in developing more effective methods for learning discriminative features from both sequence and structure modalities. Although existing DL models can automatically extract features, the complex semantics between both sequence and structure information of AMPs are still not fully exploited for representation learning.

To address these issues, this study proposes a dual diffusion model-based framework for AMPs identification, in which various physicochemical properties of AMPs are made full use by two types of contrastive learning strategies. Specifically, the peptide sequence features are derived from six methods, capturing multi-perspective sequential patterns. Simultaneously, amino acid pair interactions are introduced for structure modeling, providing a more enriched representation of the peptides' functional properties. Moreover, we design a dual diffusion process, which further refines the initial sequence and structure representations to obtain more discriminative features. For the sequence modality, we apply a Transformer-based denoising strategy, leveraging the attention mechanism to capture long-range dependencies within peptide sequences. The Transformer excels at denoising noisy or incomplete sequences, making it ideal for capturing complex relationships between amino acids in AMPs (Cao *et al.* 2023). As for the structure modality, we use a Graph Attention Network (GAT) for denoising. Given the graph-like nature of protein structures, where amino acids interact non-linearly, GATs effectively capture both local and global dependencies by focusing on the most relevant amino acids during graph representation learning (Chen *et al.* 2024b), which enhances the structural representations. Subsequently, we use single-modal contrastive learning to enhance the model's ability to distinguish between AMPs and non-AMPs, while the dual-modal contrastive learning strategy aligns the sequence and structure representations. Finally, these processed features are fused via a cross-attention mechanism and fed into a MLP classifier for

predicting AMPs activity. In summary, the main contributions of this paper are listed as follows:

- A multi-head attention mechanism is introduced, which effectively integrates local biological features with global protein language model (ESM2) embeddings for sequence data, establishing a complementary relationship.
- We propose a dual diffusion model-based representation learning framework that captures deeper biological semantics from initial features, addressing the issue of insufficient representation learning for AMPs classification.
- We also introduce two contrastive learning strategies to optimize feature representations and achieve cross-modality alignment.
- Comprehensive experimental results demonstrate the effectiveness of our approach, establishing it as a feasible solution to AMPs identification.

2 Materials and methods

This section presents the details of the proposed framework, which introduces a dual diffusion representation learning for binary-label classification of AMPs. Specifically, the proposed framework consists of three modules including the multi-view feature construction module, the Dual-modality representation learning module, and the classification module. The overall flowchart of the model is shown in Fig. 1. The detailed description of each module is presented as follows.

2.1 Multi-view feature construction

In this module, we process peptide sequences and structures through a comprehensive feature extraction pipeline. First, FASTA-formatted peptide sequences are encoded using six distinct feature extraction/representation methods to generate initial sequence representations. Concurrently, for peptides with available PDB-formatted tertiary structures, we extract atomic-level node features and edge information to construct initial graph representations. These processed sequence and structure features serve as optimized inputs for the subsequent dual-modality representation learning module.

2.1.1 Sequence-based initial feature construction

Peptide sequences are composed of amino acids, each of which is represented by a single letter. Since AMPs are generally short peptides, the maximum sequence length is set to 100. Sequences shorter than 100 residues are padded with zeros to maintain dimensional consistency across all samples. This standardization ensures uniform processing in subsequent feature extraction and model training phases. Then, six methods are used to model peptide sequence information from various perspectives, including the BLOSUM62 matrix (Henikoff and Henikoff, 1992), PAM120 matrix (Mount 2008), hydrophobicity index (Chen *et al.* 2007), helical tendency (Huang *et al.* 2010), AAindex (Kawashima *et al.* 2008), and the protein language model ESM2 (Lin *et al.* 2022). Detailed descriptions of these methods are provided in Section 2, available as [supplementary data](#) at *Bioinformatics* [online of Supplementary Material](#). Finally, the features obtained by these six methods are

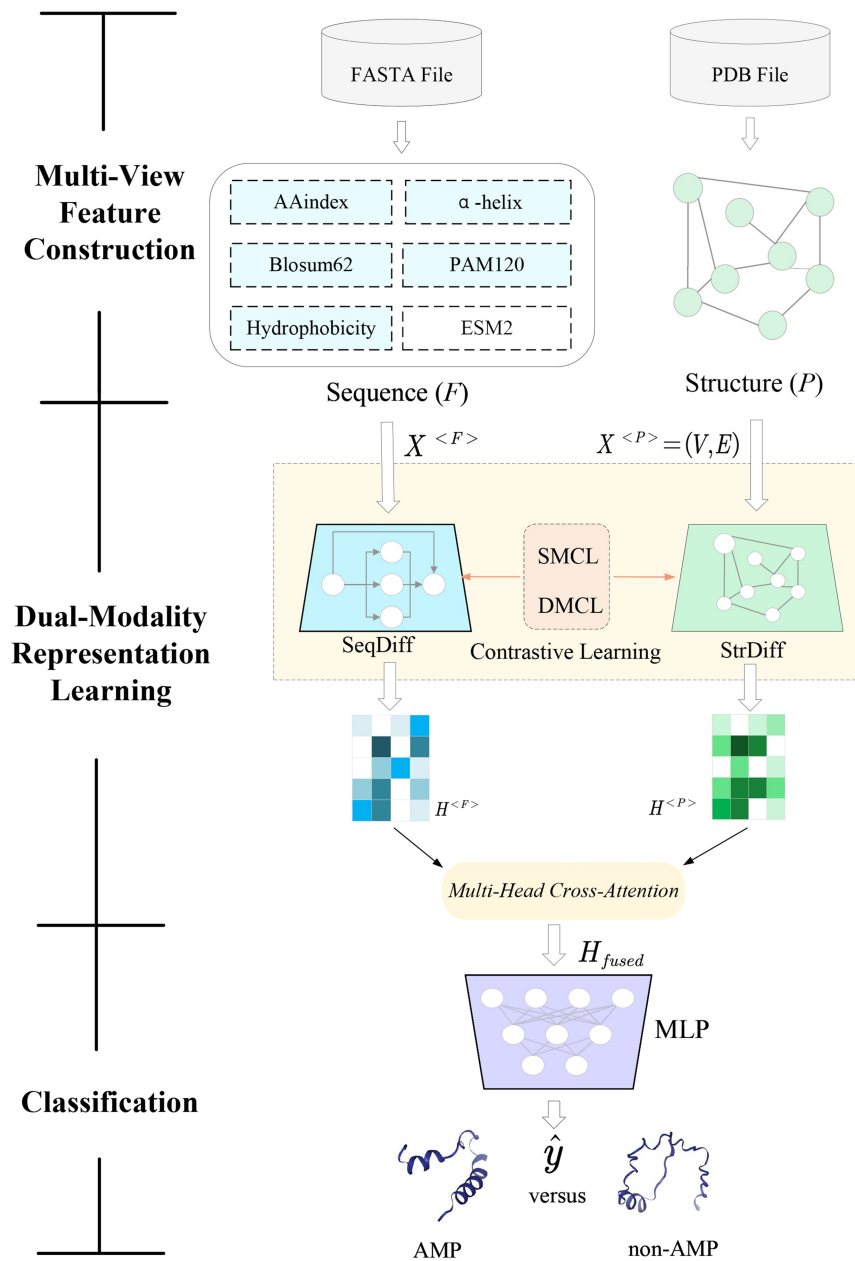


Figure 1 The schematic representation of the proposed framework. The framework consists of three modules. (i) Multi-View Feature Construction: Peptide sequences from FASTA files are encoded using multiple feature descriptors to obtain sequence features $X^{<F>}$, while peptide structures from PDB files are represented as graphs $X^{<P>} = (V, E)$. (ii) Dual-Modality Representation Learning: The sequence diffusion model (SeqDiff) and structure diffusion model (StrDiff) are used to learn latent representations $H^{<F>}$ and $H^{<P>}$ from sequence and structure modalities, respectively. Single-Modal Contrastive Learning (SMCL) and Dual-Modal Contrastive Learning (DMCL) are incorporated to enhance intra-modal and inter-modal representation consistency. Subsequently, Multi-Head Cross-Attention is applied to fuse the dual-modality representations into a unified representation H_{fused} . (iii) Classification: A Multi-Layer Perceptron (MLP) takes the fused representation as input to predict the final label $\hat{y}$, distinguishing AMPs from non-AMPs.

classified into two distinct categories: (i) Bio_feature $[100 \times 106]$ combining the first five biologically interpretable features and (ii) ESM_feature $[100 \times 2560]$ comprising data-driven embeddings with richer hidden semantics.

To effectively integrate these two heterogeneous features, we implement a three-step fusion method. First, both types of features are independently projected to an aligned 100×384 space via separate linear layers to normalize their distributions. Second, a multi-head cross-attention mechanism is applied to

capture fine-grained dependencies across the two types of features, effectively bridging the semantic gap and mitigating feature heterogeneity. Finally, the two refined features are concatenated along the feature dimension, deriving the final 100×768 sequence representation $X^{<F>}$. The above preprocessing procedure for sequence features aims to capture both local physicochemical properties and global sequence patterns, establishing a robust foundation for the subsequent representation learning.

2.1.2 Structure-based initial feature construction

In this study, the structure data in PDB format are used for deriving structure-based features. The PDB file adheres to a strict format and contains information such as atomic coordinates, chemical bonds, and crystallographic data. Building on this, the structure of the AMP is represented as a graph denoted as $X^{<P>} = (V, E)$, where V represents the set of nodes and E represents the set of edges. The process of constructing the graph with node and edge information is outlined below.

Construction of node features. In our implementation, amino acid residues of peptides are selected as nodes. The features of each node include the amino acid type, spatial coordinates of four main atoms (i.e. C_α , C, N, O), the helical propensity of the amino acid, and 45 commonly used physicochemical properties from the AAindex database. Note that these 45 AAindex features are identical to those used for sequence information, which are used as fundamental physicochemical descriptors. In the processing procedure for structure information, the AAindex features are further integrated/aggregated with topological information via GAT-based message passing strategy to derive more meaningful structure-based representations. Thereafter, the node feature matrix $\mathbf{h} \in \mathbb{R}^{N \times 78}$ is constructed. These features encompass residue type, spatial position, and biochemical properties, which are essential for understanding the functional and structure characteristics of peptides. Moreover, these features allow the model to capture the complex relationships between nodes.

Construction of edge features. First, a residue contact map of each peptide is constructed based on the coordinates of the C_α atoms. There are two conditions for establishing an edge: (i) the Euclidean distance between the C_α atoms of any two residues is less than a predefined threshold of 5 Å and (ii) the sequence distance between the pair of amino acids is less than or equal to 1. If either of these conditions is satisfied, the pair of two residues is considered to be in close proximity, and an edge is constructed between these two amino acids.

Then, the edge features incorporate multiple biologically relevant descriptors which are divided into three categories: (i) Basic geometric information including C_α - C_α distances (1D) and relative orientation vectors (3D), (ii) Residue interaction potentials comprising Distance-dependent Pairwise Statistical (DPSP) (Zhao and Xu 2012) potential (1D) and statistical pairwise interaction potentials (8D), and (iii) Sequence-derived features (2D) capturing both sequential adjacency and relative sequence positions to preserve primary structure information. These multidimensional edge attributes are collectively represented as a feature matrix $\mathbf{E}_{attr} \in \mathbb{R}^{E \times 15}$, integrating spatial, sequential, and physicochemical prior knowledge.

This multi-granular graph representation encodes three distinct levels of structure information, i.e. atomic-level (geometric coordinates), residue-level (interaction potentials), and domain-level (physicochemical patterns), and the final representation is denoted as $X^{<P>}$.

2.2 Dual-modality representation learning

In this module, more meaningful representations are derived based on the initial representations of peptide sequences and structures. Specifically, the sequence diffusion model is utilized to generate embedding information $X^{<F>}$ for the peptide sequences,

while the structure diffusion model is used to obtain the 3D spatial embeddings $X^{<P>}$ for the peptide structure. Moreover, to ensure semantic alignment between the peptide sequence and structure data, two contrastive learning approaches, i.e. Single-Modal Contrastive Learning (SMCL) and Dual-Modal Contrastive Learning (DMCL), are used to further refine the embedding representations. The final embeddings for the sequence and structure are obtained, which are denoted as $H^{<F>}$ and $H^{<P>}$, respectively.

2.2.1 Sequence feature representation optimization

We use a Transformer-based diffusion model to further perform representation learning on the peptide sequences. Due to the limit of space, the implementation process of the model is illustrated in the Fig. S1, available as [supplementary data](#) at *Bioinformatics* online.

Forward noise-adding process. Diffusion model's core capability is the restoration of the true data distribution from noise. In this study, this concept is utilized to optimize the feature learning. Through leveraging the diffusion model's ability to recover data at different noise levels, we aim to optimize feature representations by removing redundant information, mitigating interference, and learning the distribution of the initial features $X^{<F>}$. Gaussian noise is added to the input features, and noisy features are generated through the forward diffusion process:

$$X_t^{<F>} = \sqrt{\alpha_t} X_0^{<F>} + \sqrt{1 - \alpha_t} \epsilon \quad (1)$$

where $X_0^{<F>}$ denotes the initial concatenated features, t denotes the time step which controls the noise intensity, ϵ represents Gaussian noise, and $\alpha_t = \prod_{i=1}^t (1 - \beta_i)$ with $\{\beta_i\}_{i=1}^T$ denoting the noise schedule. In this study, a linear noise schedule is adopted, where β_t increases linearly from 0.0001 to 0.02 over $T = 1000$ steps. The values of α_t are precomputed and used during training to sample noisy features at arbitrary time steps.

Transformer-based sequence feature denoising. In the reverse process of diffusion model, the Transformer framework is used to capture the global dependencies within the sequence, by which the optimized representations $H^{<F>}$ are obtained by gradually predicting the noise ϵ step by step.

First, noisy feature embeddings are enhanced with sinusoidal positional encoding, enabling explicit preservation of sequence ordering relationships. Second, learnable temporal embeddings are introduced to encode noise-level information, allowing dynamic feature adaptation across different denoising steps. The main architecture consists of N stacked Transformer blocks ($N=6$ in our implementation), each of which includes multi-head self-attention (8 heads used in this study), residual connection with layer normalization (*Add & Norm*), position-wise feed-forward network (FFN), and additional *Add & Norm* operation. Specifically, the 8-head self-attention operates on the 768D token representations, where each head attends to a 96D subspace to capture intra-sequence dependencies. The process is guided by the denoising loss as Equation 2 to refine the feature representations.

$$\mathcal{L}_{\text{denoise}}^{<F>} = \|\epsilon - \epsilon_\theta(\mathbf{X}_t^{<F>}, t)\|_2^2 \quad (2)$$

Sequence modality contrastive learning. Due to the significant difference in sequence data between AMPs and non-AMPs,

Single-Modal Contrastive Learning (SMCL) is used to enhance the representation learning from the perspective of sequence modality. Furthermore, this approach is able to effectively mitigate the overfitting issue induced by limited training data, and to enhance the model's generalization capability as well.

Specifically, for each mini batch, each peptide i is treated as an anchor sample. Given peptide i , peptides having the same label as peptide i are regarded as positive samples, while other peptides with the opposite class are treated as negative samples. By maximizing the similarity of positive sample pairs and minimizing the similarity of negative sample pairs (Zheng et al. 2021), the model learns more discriminative representations for AMP identification. The loss function for the sequence modality contrastive learning is computed as follows:

$$\mathcal{L}_{\text{SMCL}}^{<F>} = -\frac{1}{N} \sum_{i=1}^N \sum_{j=1, j \neq i}^N \mathbf{1}_{y_i=y_j} \cdot \frac{\exp\left(\cos(H_i^{<F>}, H_j^{<F>})/\tau\right)}{\sum_{k=1, k \neq i}^N \exp\left(\cos(H_i^{<F>}, H_k^{<F>})/\tau\right)} \quad (3)$$

where $H_i^{<F>}$ is the representation of the i th sample, y_i denotes the label of peptide i . The indicator function $\mathbf{1}_{y_i=y_j}$ takes the value 1 when peptides i and j belong to the same class, and 0 otherwise. The function $\cos(\cdot, \cdot)$ represents the cosine similarity, τ is a temperature coefficient, and N is the size of mini batch.

Finally, the obtained features are extracted for a global representation through average pooling, followed by the application of sequence modality contrastive learning to assist in optimizing the feature representation. The obtained optimized representation is denoted by $H^{<F>}$ for the sequence modality.

2.2.2 Structure feature representation optimization

Building upon the graph-based representation of peptide structure features, the diffusion model is used to optimize the structure features of peptides, in which Graph Attention Network (GAT) is introduced as fundamental blocks. Through a process of structure noise perturbation and recovery, complex semantics among nodes (i.e. residues) are captured, and the representations for structure modality are improved accordingly. Due to the limit of space, the model architecture is shown in the Fig. S2, available as [supplementary data](#) at Bioinformatics online.

Forward noise-adding process. The structure diffusion model also includes two processes: forward noise-adding and reverse denoising. In the forward process, Gaussian noise is progressively added to the features according to the time steps, which is calculated as follows:

$$X_t^{<P>} = \sqrt{\alpha_t} X_0^{<P>} + \sqrt{1 - \alpha_t} \epsilon \quad (4)$$

Multi-scale feature encoder. To address the heterogeneous characteristics of node features, which are derived from sequences, spatial structures, and physicochemical properties, we introduce a multi-scale feature encoder that hierarchically extracts local, intermediate, and global contexts, thereby enhancing feature discriminability. By processing features at different scales, the encoder allows the diffusion model to consider the structural prior knowledge in the denoising process. This hierarchical

approach preserves fine-grained features and simultaneously captures the broader context, thereby improving the model's ability to distinguish important structural motifs from noise.

Given input node features $\mathbf{h} \in \mathbb{R}^{N \times 78}$, we first project them into a latent space $\mathbf{h}_0 \in \mathbb{R}^{N \times D}$ via a linear transformation to enrich the representational capacity. Then, the projected node features are processed through encoders at different scales, by using linear transformations to sequentially capture local, intermediate, and global semantics. Finally, the node features obtained from different scales are concatenated, resulting in a multi-scale fused node feature representation, which is denoted by h_t .

GAT-based structure feature denoising. To account for the interdependencies among residues in peptide structures, we use a GAT-based framework for structure feature denoising. The GAT leverages a multi-head attention mechanism to dynamically weight and aggregate features from neighboring nodes, enhancing contextual representations and yielding refined node embeddings (Chen et al. 2024b). Specifically, the initial hidden dimension is 256. For each GAT layer, eight attention heads are used, each of which derives a 32D hidden representation. And the eight 32D vectors are concatenated to recover a 256D node representation. Edge features are incorporated during attention computation to strengthen structural encoding. A residual connection is applied to combine the aggregated multi-head features with the original noisy input, mitigating gradient issues and preserving feature information. Finally, a global attention layer is used to capture long-range dependencies. The model is trained with a denoising objective formulated as follows:

$$\mathcal{L}_{\text{denoise}}^{<P>} = \|\epsilon - \epsilon_\theta(\mathbf{X}_t^{<P>}, t)\|_2^2 \quad (5)$$

Structure modality contrastive learning. Similar to the sequence modality, Single-Modal Contrastive Learning (SMCL) is also used to characterize the significant differences between AMPs and non-AMPs for structure data. The process of structure modality SMCL follows the same principle as the sequence modality. Specifically, the contrastive learning loss function for the structure modality is defined as follows:

$$\mathcal{L}_{\text{SMCL}}^{<P>} = -\frac{1}{N} \sum_{i=1}^N \sum_{j=1, j \neq i}^N \mathbf{1}_{y_i=y_j} \cdot \frac{\exp\left(\cos(H_i^{<P>}, H_j^{<P>})/\tau\right)}{\sum_{k=1, k \neq i}^N \exp\left(\cos(H_i^{<P>}, H_k^{<P>})/\tau\right)} \quad (6)$$

Finally, we obtain the optimal graph-level feature representation, which is denoted by $H^{<P>}$.

2.2.3 Dual-modal contrastive learning

While SMCL focuses on learning class-discriminative representations within each modality, Dual-Modal Contrastive Learning (DMCL) is used to explicitly align the sequence and structure representations in a shared semantic space, aiming to bridge the modality gap. Through maximizing cross-modal agreement via an InfoNCE objective function, the DMCL is able to jointly optimize sequence and structure representations. Due to the space limit, the schematic illustration is shown in Fig. S3, available as [supplementary data](#) at Bioinformatics online.

Specifically, the sequence (or structure) of a peptide is used as the anchor, with its corresponding structure (or sequence) as the positive sample, while the structures (or sequences) of other peptides are used as negative samples. The objective is that the mutual information (MI) between the anchor and positive samples is maximized, while the MI between the anchor and negative samples is minimized. In this study, the InfoNCE loss (Wang et al. 2024) is used to compute the contrastive learning loss between modalities, which is formally defined as follows:

$$\mathcal{L}_{\text{DMCL}} = -\frac{1}{2N} \sum_{i=1}^N \left[\log \frac{\exp(\cos(H_i^{<F>}, H_i^{<P>})/\tau)}{\sum_{j=1}^N \exp(\cos(H_i^{<F>}, H_j^{<P>})/\tau)} + \log \frac{\exp(\cos(H_i^{<P>}, H_i^{<F>})/\tau)}{\sum_{j=1}^N \exp(\cos(H_i^{<P>}, H_j^{<F>})/\tau)} \right] \quad (7)$$

where $H_i^{<F>}$ and $H_i^{<P>}$ denote the sequence and structure representations of the i th peptide, respectively. The two parts in the square bracket of right hand of Equation (7) correspond to sequence-to-structure and structure-to-sequence contrastive objectives, and the factor $1/2$ averages the loss over the two directions.

2.2.4 Training objective

Overall, the loss function for training the module of Dual-Modality Representation Learning is defined as follows:

$$\mathcal{L}_{\text{DMRL}} = \lambda_3 \left(\lambda_1 \mathcal{L}_{\text{denoise}}^{<F>} + (1 - \lambda_1) \mathcal{L}_{\text{SMCL}}^{<F>} \right) + \lambda_4 \left(\lambda_2 \mathcal{L}_{\text{denoise}}^{<P>} + (1 - \lambda_2) \mathcal{L}_{\text{SMCL}}^{<P>} \right) + \lambda_5 \mathcal{L}_{\text{DMCL}} \quad (8)$$

By using the training process, we obtain the representations for sequence and structure through iterative updates, and the representations are denoted as $H^{<F>}$ and $H^{<P>}$, respectively.

Moreover, we use a cross-attention mechanism and a gating mechanism to effectively integrate the sequence and structure features. Firstly, the representation $H^{<F>}$ and $H^{<P>}$ are projected into a shared space and fused via 8-head cross-attention mechanism to capture modality interactions. Specifically, the final representation is a 512D vector, i.e. for each attention head, a 64D representation is derived. Subsequently, a gating mechanism dynamically adjusts the contribution of each modality. Finally, $H^{<F>}$ and $H^{<P>}$ are weighted by learnable parameters and concatenated to form the fused representation H_{fused} , which is used for the downstream classification task.

2.3 Classification module

In this study, a Multi-Layer Perceptron (MLP) is used as the classifier for AMPs classification. Specifically, the representations obtained from the dual-modality feature learning module are utilized as input, and the output of the classification represents whether the input peptide exhibits antimicrobial activity or not.

2.3.1 MLP model architecture

In the three-layer MLP, two fully connected layers with ReLU activation and Dropout are used to learn enriched representations. The final output layer performs linear transformation followed

by a sigmoid activation function to derive the probability of being AMP, completing the binary classification task accordingly.

2.3.2 Loss function

The cross-entropy loss function is used to evaluate the discrepancy between the predicted outputs and the true labels. The formula for the loss function is defined as follows:

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (9)$$

where y_i represents the true label of the i th sample, and $\hat{y}_i$ is the model's predicted label for the i th sample.

In this study, the AdamW optimizer is used to minimize the loss function, resulting in an optimized framework for AMPs classification.

3 Experimental setup

3.1 Dataset

In experiments, the positive samples or AMPs are collected from several AMP databases, including CAMP, DBAASP, DRAMP, APD3, and SATPDB (Wang, 2023). After applying a preprocessing procedure (including removing duplicates, excluding peptides containing non-standard amino acids, and filtering out sequences longer than 100 residues), we obtain a dataset of 10 057 AMPs. This dataset is split into training, validation, and testing subsets in the ratio of 8:1:1. For negative samples, we collect the equal number of non-AMPs from the UniProt database to pair with the AMPs dataset for contrastive learning. It is worth noting that for about 67.4% of AMPs sequences without available validated structures, we utilize ESMfold (Lin et al. 2022) to predict their corresponding structure information to complete the structural modality.

3.2 Hyperparameter settings

In experiments, the proposed model is implemented based on PyTorch (<https://pytorch.org/>). Generally, the number of epochs and batch size are set as 100 and 32, respectively. The AdamW optimizer is used to train the whole model with early stopping (patience=10). Note that AdamW is particularly effective for training multi-stage framework where stable convergence and generalization are prioritized. This obtained architecture configuration balances model complexity, computational efficiency, and generalization capability while effectively capturing sequence-structure feature representation in AMPs.

Note that the whole model contains many hyperparameters, and due to space limit, detailed setting of hyperparameters is placed in Section 3, available as [supplementary data at Bioinformatics online of Supplementary Material](#). For all hyperparameters, three methods are used to obtain the settings. Firstly, some of them are set as the default values in PyTorch library, including weight decay and dropout rate. Secondly, some of them are set by following settings in previous studies, including number of layers, number of attention heads, hidden dimension, layer dimension, optimizer, batch size and learning rate. For two framework-level hyperparameters,

i.e. the number of timesteps and the noise schedule for diffusion model, the experiment of hyperparameter sensitivity analysis is conducted to derive the final setting. The results of hyperparameter analysis are presented in [Section 9](#), available as [supplementary data](#) at *Bioinformatics online of Supplementary Material*.

3.3 Baselines and evaluation metrics

In order to evaluate the proposed model, we select nine representative methods as baselines, including iAMP-Attenpred ([Xing et al. 2023](#)), AI4AMP ([Lin et al. 2021](#)), AMP-BERT ([Lee et al. 2023](#)), sAMPpred-GAT ([Yan et al. 2023](#)), amPEPpy ([Lawrence et al. 2021](#)), Ampir ([Fingerhut et al. 2020](#)), iAMP-DL ([Nguyen et al. 2024](#)), deepAMPNet ([Zhao et al. 2024](#)), and TP-LMMSG ([Chen et al. 2024a](#)). The brief introduction of these baselines is presented in [Section 4](#), available as [supplementary data](#) at *Bioinformatics online of Supplementary Material*.

In experiments, we evaluate the proposed method against these baselines according to five widely adopted metrics: accuracy (ACC), F1 score, sensitivity (Sn), specificity (Sp), and the area under the ROC curve (AUC).

To ensure a fair and reproducible comparison, the same data splits are used for all models. Additionally, baseline models are run with the hyperparameters tuned according to the original papers. Due to the randomness in data splitting and model initialization, the results of the baseline methods may slightly differ from those reported in the original papers.

4 Results and analysis

4.1 Overall comparison

We first compare the overall performance of our model with selected baselines according to various evaluation metrics, and the experimental results are summarized in [Table S2](#), available as [supplementary data](#) at *Bioinformatics online*.

From results in [Table S2](#), available as [supplementary data](#) at *Bioinformatics online*, we can find that the proposed method consistently achieves superior performance across all metrics compared with the baseline methods. In addition, DL-based baselines (iAMP-Attenpred, AI4AMP, AMP-BERT, and sAMPpred-GAT) generally perform better than traditional machine learning-based methods (amPEPpy and Ampir), indicating the superior capability of DL for deriving more significant features. Due to space limit, the more detailed results and analysis are provided in [Section 4.2](#), available as [supplementary data](#) at *Bioinformatics online of the Supplementary Material*.

4.2 Ablation study

To validate the contribution of main components in our model to the performance improvement, we compare our model with its seven different variants. Specifically, the main components include single-modal contrastive learning, dual-modal contrastive learning, whole contrastive learning, multi-scale encoder, sequence diffusion model, structure diffusion model, and ESMfold for deriving predicted structures. In addition, we also evaluate the impact of predicted structures. The experimental results are presented in the [Table S3](#), available as [supplementary data](#) at *Bioinformatics online*.

Generally, the results indicate that seven components are all essential for improving the prediction performance on AMPs. Moreover, both sequence diffusion model and structure diffusion model make more contributions than other components, suggesting that both components are really important for deriving more significant representations. Owing to space limit, the detailed results and analysis are provided in [Section 6](#), available as [supplementary data](#) at *Bioinformatics online of Supplementary Material*.

4.3 Evaluation of sequence diffusion representation learning

To further investigate the effectiveness of the sequence diffusion representation learning model for optimizing sequence features, we qualitatively compare the distribution of obtained representations for AMPs and non-AMPs, the results of which are shown in the [Fig. S4](#), available as [supplementary data](#) at *Bioinformatics online*. Additionally, three representative features are selected for quantitative comparison and statistical significance analysis.

[Figure S4](#), available as [supplementary data](#) at *Bioinformatics online* illustrates the feature distributions of AMPs and non-AMPs, providing critical insights into the class-specific feature patterns. From [Fig. S4](#), available as [supplementary data](#) at *Bioinformatics online*, we can find that the sequence diffusion representation learning model really derives the discriminative features, which are able to basically characterize the difference between AMPs and non-AMPs. Specifically, the mean difference is about 1.1 across all features, indicating the significant difference between different classes (i.e. AMPs and non-AMPs). Moreover, the statistical test results exhibit P -value $< .001$ and Cohen's $d > 1.5$, suggesting the significance of the comparison results. Therefore, the sequence diffusion representation learning module can effectively capture discriminative semantics which is critical for AMPs classification.

To further elucidate the effectiveness of sequence diffusion representation learning, kernel density estimation ([Parzen 1962](#)) is applied to the features with the top three largest mean difference, and violin plots are used to visualize the distributions of three features (shown in the [Fig. S5](#), available as [supplementary data](#) at *Bioinformatics online*). As shown in [Fig. S5](#), available as [supplementary data](#) at *Bioinformatics online*, the feature F341 exhibits almost no overlap in the distributions of two classes (mean difference = 3.13, P -value = 2×10^{-10} , and Cohen's $d = 4.41$), indicating the significant difference on feature F341 between AMPs and non-AMPs. The similar patterns can be observed for feature F229 (mean difference = 3.09, P -value = 5×10^{-10} , Cohen's $d = 4.15$) and feature F34 (mean difference = 3.01, P -value = 8×10^{-10} , Cohen's $d = 4.27$). These findings demonstrate that the features obtained by the sequence diffusion representation learning module can effectively discriminate the AMPs and non-AMPs, improving the performance of AMPs classification accordingly.

4.4 Evaluation of structure diffusion representation learning

To further evaluate the effectiveness of the structure diffusion representation learning model, we also compare the distributions of average features derived by this model for AMPs and

non-AMPs, which are shown in Fig. S6, available as [supplementary data](#) at *Bioinformatics* online.

Figure S6, available as [supplementary data](#) at *Bioinformatics* online illustrates that the structure diffusion representation learning model can derive discriminative features that distinctly characterize both AMPs and non-AMPs. Specifically, the average feature difference is approximately 1.63, which is greater than that obtained by sequence diffusion model (with mean difference about 1.1). Moreover, the statistical analysis confirms the significance of this distinction, with the P -value $< .001$ and the mean Cohen's d value > 1.5 . These results demonstrate that the structure diffusion representation learning model can effectively capture the potential semantics between AMPs and non-AMPs, which are essential for accurate AMP classification.

4.5 Visualization of representation learning

To gain a deeper understanding of the role of each component in the dual-modality feature learning module, we utilize t-SNE to visualize obtained features at different stages, which can intuitively demonstrate the changes of feature distribution across different stages. The results are shown in Fig. 2. Note that visualization highlights the difference of obtained features between AMPs and non-AMPs throughout the training process.

For sequence features, the representations before and after sequence diffusion model are shown in Fig. 2a and b, respectively. While the sequence diffusion model provides some improvement, there is still certain overlap between AMPs and non-AMPs representations. For structure features, the representations before and

after structure diffusion model are shown in Fig. 2c and d. It is evident that the structure diffusion module effectively separates AMPs and non-AMPs samples, with post-training feature representations demonstrating a more concentrated and distinct distribution. This indicates that the structure diffusion module effectively captures complex patterns, enhancing the discriminative capabilities of the features. Note that this observation is consistent with the results in Sections 4.3 and 4.4, i.e. the mean feature difference derived by sequence diffusion model between AMPs and non-AMPs is about 1.1 while the counterpart derived by structure diffusion model is 1.63. Figure 2e illustrates the fused representation of sequence and structure information. From the t-SNE visualization, it is clear that the separation between AMPs and non-AMPs samples is further improved, indicating that the effective fusion of sequence and structure representations provides a more discriminative feature space for AMPs classification tasks. Figure 2f shows the feature representations after passing through the classification model. Compared to Fig. 2e, the classification model further enhances the separability of the features, achieving a more distinct partitioning of AMPs and non-AMPs samples in the 2D space. This demonstrates that, after classification model training, the features not only become more semantically precise but also allow the model to map these features into decision spaces more effectively.

In summary, the distribution of features for AMPs and non-AMPs becomes progressively clearer as the model training proceeds. This process intuitively illustrates how the initial noisy representations evolve into well-defined classification boundaries, showcasing the effectiveness of the model for deriving more meaningful features.

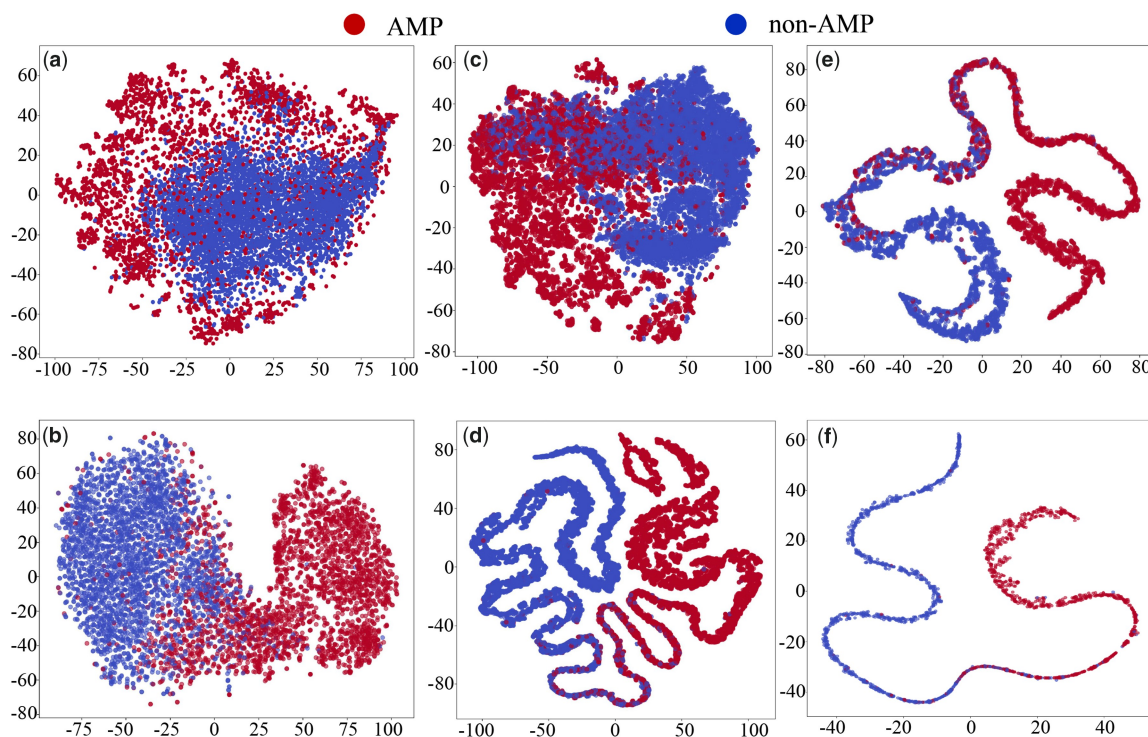


Figure 2 Visualization of learned representations at different stages. (a) Before sequence diffusion representation learning. (b) After sequence diffusion representation learning. (c) Before structure diffusion representation learning. (d) After structure diffusion representation learning. (e) After dual-modality fusion feature learning. (f) After MLP layers.

5 Discussion

Accurate identification of AMPs is the critical first step in developing novel antimicrobial therapies. However, existing AMPs prediction methods struggle with limitations in the depth of feature extraction and the ability to effectively fuse multi-modal information. In this study, we propose a dual diffusion model-based representation learning framework, which has the following key advantages. (i) Sequences and structures of peptides are used from multiple perspectives, providing high-quality initial representations. For sequence information, multiple feature descriptors are leveraged to effectively integrate local and global semantics. For structure information, three hierarchical levels of structural semantics are extensively characterized. (ii) The dual diffusion model architecture is utilized to learn more discriminative representations from both sequence and structure perspectives. (iii) Two contrastive learning strategies are introduced as well to optimize further the representations of peptides, improving the performance of AMPs identification accordingly.

Generally speaking, the proposed framework in this study holds the potential for developing alternative antimicrobial agents. The high predictive power of our framework can significantly shorten the cycle for screening potential AMPs from vast peptide libraries, helping to reduce early-stage failure rates in drug discovery and accelerate the identification of lead compounds. Furthermore, the proposed framework can be readily adapted to other functional peptide prediction tasks (e.g. anti-fungal peptides, cell-penetrating peptides), demonstrating broad potential for biomedical applications.

Despite the positive results achieved in this study, the proposed framework still has the following limitations. (i) The feature generation within the dual diffusion process remains a black box, limiting the understanding of feature evolution in representation learning. Our future work will try to address this issue through introducing explainable AI techniques. (ii) The introduction of the dual diffusion model and contrastive learning strategies increases model complexity, leading to higher demands on training time and computational resources. (iii) In our study, the non-AMPs samples are collected from UniProt and may include peptides with unannotated antimicrobial activity, introducing label noise. Furthermore, the limited training data may not fully cover the vast heterogeneity of AMPs. Future efforts will focus on expanding datasets and incorporating more diverse peptide sequences and structures to enhance model generalizability.

6 Conclusions

In this study, we have proposed a dual diffusion model-based representation learning framework for AMPs classification. Our method leverages both sequence and structure information of peptides to address the limitations of existing AMPs prediction methods. First, the multi-view feature construction module is utilized to characterize peptide sequences and structures from multiple perspectives, resulting in enriched biological embeddings. Second, the dual diffusion representation learning module is used to further optimize sequence-based features and

structure-based features. Third, the contrastive learning strategy is used to effectively enhance the model's ability to differentiate between AMPs and non-AMPs while aligning sequence and structure information. Finally, comprehensive experimental results demonstrate that our model outperforms existing methods for AMPs classification. Generally, the proposed dual diffusion model framework holds significant advantages in deep feature mining and multi-modal information fusion, which cannot only improve predictive performance but also lay a technical foundation for accelerating antimicrobial drug discovery.

Acknowledgements

Authors are grateful to the anonymous reviewers for helpful comments.

Author contributions

Wen Kong (Investigation [equal], Writing—original draft [equal]), Lingling Fu (Writing—review & editing [equal]), and Xingpeng Jiang (Conceptualization [equal], Writing—review & editing [equal])

Supplementary material

[Supplementary material](#) is available at *Bioinformatics* online.

Conflicts of interest

None declared.

Funding

This work was partially supported by the National Natural Science Foundation of China [62472192 and 62372205].

References

- Al-Omari AM, Akkam YH, Zyout A *et al.* Accelerating antimicrobial peptide design: leveraging deep learning for rapid discovery. *PLoS One* 2024;**19**:e0315477.
- Alizadehsani R, Oyelere SS, Hussain S *et al.* Explainable artificial intelligence for drug discovery and development: a comprehensive survey. *IEEE Access* 2024;**12**:35796–812.
- Brandes N, Ofer D, Peleg Y *et al.* Proteinbert: a universal deep-learning model of protein sequence and function. *Bioinformatics* 2022;**38**:2102–10.
- Cao Q, Ge C, Wang X *et al.* Designing antimicrobial peptides using deep learning and molecular dynamic simulations. *Brief Bioinform* 2023;**24**:bbad058.
- Chen N, Yu J, Liu Z *et al.* Tp-lmmsg: a peptide prediction graph neural network incorporating flexible amino acid property representation. *Brief Bioinform* 2024a;**25**:bbae308.
- Chen S, Tang Z, You L *et al.* A knowledge distillation-guided equivariant graph neural network for improving protein

- interaction site prediction performance. *Knowl Based Syst* 2024b;**300**:112209.
- Chen Y, Guarnieri MT, Vasil AI *et al.* Role of peptide hydrophobicity in the mechanism of action of α -helical antimicrobial peptides. *Antimicrob Agents Chemother* 2007;**51**:1398–406.
- Fingerhut LC, Miller DJ, Strugnell JM *et al.* ampir: an r package for fast genome-wide prediction of antimicrobial peptides. *Bioinformatics* 2020;**36**:5262–3.
- Garzon Otero D, Akbari O, Bilodeau C. Pepmnet: a hybrid deep learning model for predicting peptide properties using hierarchical graph representations. *Mol Syst Des Eng* 2025;**10**:205–18.
- Henikoff S, Henikoff JG. Amino acid substitution matrices from protein blocks. *Proc Natl Acad Sci USA* 1992;**89**:10915–9.
- Huang Y, Huang J, Chen Y. Alpha-helical cationic antimicrobial peptides: relationships of structure and function. *Protein Cell* 2010;**1**:143–52.
- Kawashima S, Pokarowski P, Pokarowska M *et al.* Aaindex: amino acid index database, progress report 2008. *Nucleic Acids Res* 2008;**36**:D202–5.
- Lawrence TJ, Carper DL, Spangler MK *et al.* ampeppy 1.0: a portable and accurate antimicrobial peptide prediction tool. *Bioinformatics* 2021;**37**:2058–60.
- Lee H, Lee S, Lee I *et al.* Amp-bert: prediction of antimicrobial peptide function based on a Bert model. *Protein Sci* 2023;**32**:e4529.
- Lin T-T, Yang L-Y, Lu I-H *et al.* Ai4amp: an antimicrobial peptide predictor using physicochemical property-based encoding method and deep learning. *mSystems* 2021;**6**:e00299–21.
- Lin Z, Akin H, Rao R *et al.* Language models of protein sequences at the scale of evolution enable accurate structure prediction. bioRxiv, 2022, preprint: not peer reviewed.
- Mount DW. Comparison of the pam and blosum amino acid substitution matrices. *CSH Protoc* 2008;**2008**:pdb.ip59.
- Nguyen QH, Nguyen-Vo TH, Do TTT *et al.* An efficient hybrid deep learning architecture for predicting short antimicrobial peptides. *Proteomics* 2024;**24**:e2300382.
- Parzen E. On estimation of a probability density function and mode. *Ann Math Statist* 1962;**33**:1065–76.
- Ruiz Puentes P, Henao MC, Cifuentes J *et al.* Rational discovery of antimicrobial peptides by means of artificial intelligence. *Membranes (Basel)* 2022;**12**:708.
- Salam MA, Al-Amin MY, Salam MT *et al.* Antimicrobial resistance: a growing serious threat for global public health. *Healthcare* 2023;**11**:1946.
- Wang G. The antimicrobial peptide database is 20 years old: recent developments and future directions. *Protein Sci* 2023;**32**:e4778.
- Wang Y, Liang V, Yin N *et al.* Sgac: a graph neural network framework for imbalanced and structure-aware amp classification. *Brief Bioinform*, 2026;**27**:bbag038.
- Wang Z, Wang Y, Zhang W. Improving paratope and epitope prediction by multi-modal contrastive learning and interaction informativeness estimation. In: Larson K (ed.), *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence, IJCAI-24. International Joint Conferences on Artificial Intelligence Organization. Main Track.* 2024, 6053–61.
- Xing W, Zhang J, Li C *et al.* iamp-attenpred: a novel antimicrobial peptide predictor based on Bert feature extraction method and CNN-BILSTM-attention combination model. *Brief Bioinform* 2023;**25**:bbad443.
- Yan K, Lv H, Guo Y *et al.* TPPRED-ATMV: therapeutic peptide prediction by adaptive multi-view tensor learning model. *Bioinformatics* 2022;**38**:2712–8.
- Yan K, Lv H, Guo Y *et al.* sampred-gat: prediction of antimicrobial peptide by graph attention network and predicted peptide structure. *Bioinformatics* 2023;**39**:btac715.
- Zhao F, Xu J. A position-specific distance-dependent statistical potential for protein structure and functional study. *Structure* 2012;**20**:1118–26.
- Zhao F, Qiu J, Xiang D *et al.* deepampnet: a novel antimicrobial peptide predictor employing alphafold2 predicted structures and a bi-directional long short-term memory protein language model. *PeerJ* 2024;**12**:e17729.
- Zheng M, Wang F, You S *et al.* Weakly supervised contrastive learning. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision.* 2021, 10042–51.