

Dual-Topology Learning with Adaptive Anchors for Multi-View Clustering

Chenglong Zhang¹, Chao Zhang¹, Junhao Zhang¹, Junyi Guan², Xianzhong Zhou¹,
Bo Wang^{1*} and Huaxiong Li^{1*}

¹Nanjing University, Nanjing, China

²Hangzhou Normal University, Hangzhou, China

clzhang123@163.com, czhang@nju.edu.cn, junhao.zhang@smail.nju.edu.cn,
jonnyguan73@163.com, zhouxz@nju.edu.cn, bowangsm@nju.edu.cn, huaxiongli@nju.edu.cn

Abstract

As a prominent paradigm for large-scale unsupervised learning, anchor-based multi-view clustering aims to reveal the latent structures across heterogeneous data representations with high efficiency. Despite achieving some progress, existing methods typically suffer from the following two limitations. Firstly, they rely on pre-constructed anchors or rigid constraints (e.g., orthogonality), whereas the intrinsic topological correlations among anchors are completely discarded. Secondly, most existing methods either perform clustering directly on the bipartite graph or treat different views equally, thereby failing to capture the global structural information. To this end, the Dual-Topology learning with adaptive Anchors (DOTA) is proposed, which not only learns the sample-anchor relationship (i.e., bipartite graph) but also preserves the topology structure among anchors, significantly enhancing the discriminability of learned representation while preserving the underlying data manifold. By integrating an adaptive view-weighting strategy to balance view contributions, DOTA derives discriminative global sample embeddings by propagating spectral information through the bipartite graph. Extensive experiments on benchmark datasets demonstrate the superiority of DOTA.

1 Introduction

With the rapid advancement of information technologies, data collected from heterogeneous sources has emerged extensively, referred to as multi-view data [Zhang *et al.*, 2024b; Wang *et al.*, 2025; Fan *et al.*, 2025; Kou *et al.*, 2025a]. Accordingly, Multi-View Clustering (MVC), which aims to leverage complementary information from multiple views to uncover the intrinsic data structures, has become a pivotal unsupervised learning paradigm. Based on the sparse representation assumption that a sample can be approximately reconstructed by its neighbors, many graph-based multi-view clustering methods have been proposed [Wang *et al.*, 2019; Jiang *et al.*, 2025b; Liu *et al.*, 2025]. However, the prohibitive

computational burden of constructing and decomposing the n -order graph incurs the complexity of $(\mathcal{O}(n^2))$ and $(\mathcal{O}(n^3))$, severely restricting their applications on large-scale MVC tasks [Xu *et al.*, 2024; Jiang *et al.*, 2025a; Lu *et al.*, 2025; Zhao *et al.*, 2025; Jiang *et al.*, 2026; Fan *et al.*, 2026].

To this end, anchor-based graph learning has been widely adopted in MVC, effectively reducing the computational cost by constructing the bipartite graph between m anchors and n samples ($m \ll n$). Depending on the anchor generation mechanism, existing methods can be mainly divided into two categories. Specifically, the first category leverages a two-stage strategy, where anchors are initially generated (sampling or heuristic) and then employed as static points for graph construction. For example, UDBGL [Fang *et al.*, 2023] utilizes K-means on concatenated features to obtain anchors, and AEVC [Liu *et al.*, 2024b] proposes to refine anchors under the guidance of a pre-constructed graph. However, these methods completely decouple anchor generation from the subsequent clustering process, rendering the performance highly dependent on the initial anchor selection and making the results susceptible to noise and selection bias [Zhang *et al.*, 2024a; Tang *et al.*, 2025; Liu *et al.*, 2024a].

To fully exploit the knowledge interaction between anchor learning and graph construction, the second category dynamically updates anchors during the optimization process. Representative methods include RCAGL [Liu *et al.*, 2024c], AGLDR [Chen *et al.*, 2025b] and ALPC [Chen *et al.*, 2025a]. Commonly, these methods impose orthogonal constraints to enhance anchor diversity and reduce redundancy, whereas the intrinsic correlations among anchors are neglected, failing to capture the complex similarity structure of real-world data. To alleviate this issue, [Zhang *et al.*, 2024a] proposed CAMVC, which enforces explicit cluster structure via a pre-defined cluster indicator matrix. Recently, HRAL [Zeng *et al.*, 2025] employs a hypergraph derived from bipartite graph as a regularizer to preserve the relation among anchors. Despite making some progress, these methods primarily concentrate on imposing specific structural constraints (e.g., orthogonality or cluster priors) to guide anchor learning, whereas the inflexible assumptions often lead to significant distribution distortion as the intrinsic correlations among anchors are completely overlooked. Furthermore, existing methods typically integrate multi-view information via simple summation

*The corresponding authors.

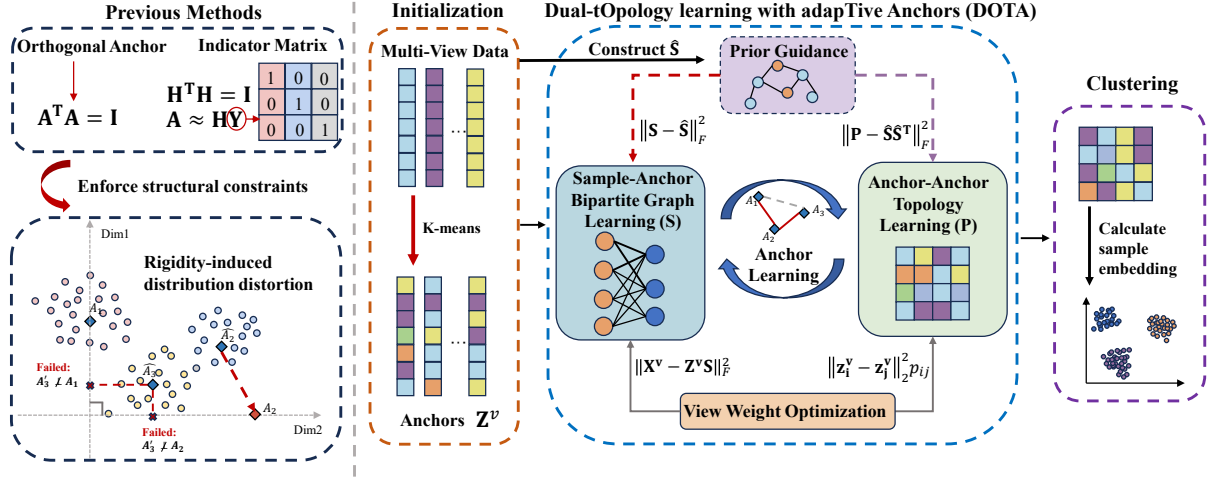


Figure 1: Schematic illustration of DOTA. Concretely, DOTA initializes the anchors by performing K-means on the concatenated multi-view features. Subsequently, the sample-anchor bipartite graph $\mathbf{S}$ and the anchor-anchor topology $\mathbf{P}$ are jointly optimized within a unified framework, where adaptive view-weights α_v are incorporated to balance the contribution of different views. Therefore, the sample-anchor as well as anchor-anchor relations can be mutually reinforced. Finally, the global sample embeddings are derived by propagating spectral information from the anchor space, effectively capturing the intrinsic structure to facilitate the ultimate clustering.

and conduct clustering directly on the bipartite graph, which not only neglects view discrepancies but also fails to capture the discriminative global structure [Jiang *et al.*, 2023; Liang *et al.*, 2025; Kou *et al.*, 2025b; Huang *et al.*, 2025].

Motivated by the aforementioned issues, we propose a Dual-topology learning with adaptive Anchors (DOTA) framework for Multi-View Clustering. The basic framework of DOTA is illustrated in Fig. 1, and the main contributions of this paper are summarized as follows:

- A unified multi-view clustering framework is proposed by seamlessly integrating adaptive anchor optimization with topology structure learning, whereby the mutual reinforcement between anchor learning and graph construction is leveraged to strengthen the adaptability to complex datasets in practical applications.
- A novel topology regularization mechanism is designed to uncover latent correlations by simultaneously optimizing the sample-anchor bipartite graph and anchor-level similarities, thereby avoiding the adverse effects of rigid constraints and effectively preserving the intrinsic topological structure among anchors.
- An adaptive supervision and fusion strategy is introduced, which not only incorporates data distribution priors to guide the learning process but also adaptively optimizes view weights to discriminate different views, ensuring the robustness of global structure information against data noise and initialization bias.

2 Methodology

Throughout this paper, matrices and vectors are denoted by bold uppercase and bold lowercase letters, respectively. Given a matrix $\mathbf{M}$, $\mathbf{m}_i$ represents its i -th row, while $\text{Tr}(\mathbf{M})$ and $\|\mathbf{M}\|_F = \sqrt{\text{Tr}(\mathbf{M}^T \mathbf{M})}$ denote the Trace and Frobenius norm. The summary of key notations is provided in Table 1.

Notation	Description
n, m	Number of samples and anchors
c, V	Number of clusters and views
d_v	Feature dimension of the v -th view
$d = \sum_{v=1}^V d_v$	Total dimension of all views
$\mathbf{x}_i^v \in \mathbb{R}^{d_v \times 1}$	The i -th sample in the v -th view
$\mathbf{z}_j^v \in \mathbb{R}^{d_v \times 1}$	The j -th anchor in the v -th view
$\mathbf{X}^v \in \mathbb{R}^{d_v \times n}$	Sample matrix of the v -th view
$\mathbf{Z}^v \in \mathbb{R}^{d_v \times m}$	Anchor matrix of the v -th view
$\mathbf{S} \in \mathbb{R}^{m \times n}$	Bipartite graph matrix between samples and anchors
$\mathbf{P} \in \mathbb{R}^{m \times m}$	Similarity graph matrix among anchors
$\mathbf{1}_n \in \mathbb{R}^{n \times 1}$	All-one vector

Table 1: Description of Notations

2.1 Adaptive Bipartite Graph Construction

While graph-based methods excel at capturing intrinsic local information, they typically require constructing the n -order dense matrix to characterize similarity relationships, resulting in significant computational costs. To this end, based on the assumption that each sample can be reconstructed by a set of representative anchors, DOTA characterizes this relationship as a sample-anchor bipartite graph, formulated as follows:

$$\min_{\mathbf{Z}^v, \mathbf{S}, \mathbf{1}_n = \mathbf{1}_m, \mathbf{S} \geq 0} \|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2 + \beta \|\mathbf{S} - \hat{\mathbf{S}}\|_F^2, \quad (1)$$

where $\mathbf{Z}^v$ represents the learned anchor matrix for v -th view, $\mathbf{S}$ and $\hat{\mathbf{S}}$ denote the adaptively optimized and pre-constructed bipartite graphs, respectively. Unlike conventional two-stage strategies (e.g., sampling or clustering centers) that decouple anchor learning from graph construction, DOTA simultaneously optimizes the view-specific anchors (i.e., $\mathbf{Z}^v$) and the unified bipartite graph (i.e., $\mathbf{S}$), which not only preserves the diversity of each view but also utilizes the unified information to ensure cross-view alignment of anchors. Furthermore, the second term in Eq. (1) incorporates the pre-constructed graph

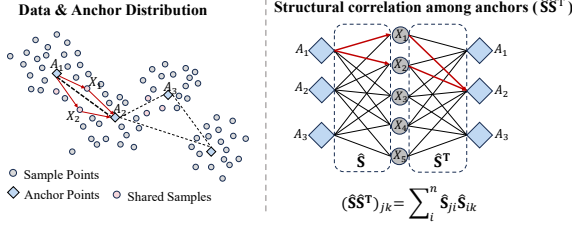


Figure 2: Mechanism of capturing prior structural correlation among anchors ($\hat{\mathbf{S}}\hat{\mathbf{S}}^T$) via shared sample connectivity.

$\hat{\mathbf{S}}$ to supervise the learning process, thereby leveraging prior knowledge to avoid trivial solutions and enhance robustness.

2.2 Prior-Driven Anchor Similarity Learning

Since anchors serve as representative prototypes that reveal the underlying topological structure, enhancing their discriminability becomes the fundamental objective in clustering tasks. Notably, existing methods primarily focus on exploring sample-anchor relationships or enforcing rigid constraints to ensure diversity, whereas the topological correlations among anchors are typically neglected. Therefore, the local information within anchor space is exploited to further preserve the underlying manifold of anchors, formulated as:

$$\min_{\mathbf{P} \mathbf{1}_m = \mathbf{1}_m, \mathbf{P} \geq 0, \text{diag}(\mathbf{P}) = 0} \sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij} + \beta \mathcal{L}(\mathbf{P}), \quad (2)$$

where $\mathbf{z}_i^v$ denotes the i -th anchor in the v -th view, p_{ij} measures the similarity between the i -th and j -th anchors, and $\mathcal{L}(\mathbf{P})$ represents the regularization term. Given that $\hat{\mathbf{S}} \in \mathbb{R}^{m \times n}$ depicts sample-anchor relationships, the prior knowledge of anchor-anchor similarity can be derived by computing $\hat{\mathbf{S}}\hat{\mathbf{S}}^T \in \mathbb{R}^{m \times m}$, where high-order similarity information can be effectively propagated through the connectivity of shared samples (as illustrated in Fig. 2). By incorporating the prior structural information, Eq. (2) can be reformulated as:

$$\min_{\mathbf{P} \mathbf{1}_m = \mathbf{1}_m, \mathbf{P} \geq 0, \text{diag}(\mathbf{P}) = 0} \sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij} + \beta \|\mathbf{P} - \hat{\mathbf{S}}\hat{\mathbf{S}}^T\|_F^2. \quad (3)$$

Unlike previous methods that rely on pre-defined constraints, DOTA adaptively explores the similarities among anchors under the guidance of prior knowledge, which facilitates an effective description of intrinsic structural information and avoids the distribution distortion caused by inflexible structural assumptions, enhancing the overall performance.

2.3 Overall Model of DOTA

To fully leverage the interaction between different structural information, DOTA proposes a unified framework that jointly optimizes the sample-anchor bipartite graph $\mathbf{S}$ and the anchor-anchor topological structure $\mathbf{P}$, so that the anchor learning and graph construction can be mutually reinforced. Additionally, considering the heterogeneous statistical properties across views, assigning uniform weights fails to distinguish the varying contributions of diverse views. Therefore, by integrating the adaptive view-weighting strategy, the overall objective function of DOTA is formulated as follows:

$$\min_{\mathbf{Z}^v, \mathbf{S}, \mathbf{P}, \alpha} \sum_{v=1}^V \alpha_v^2 \underbrace{\|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2}_{\text{Reconstruction}} + \lambda \underbrace{\sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij}}_{\text{Manifold Regularizer}} + \beta \underbrace{\|\mathbf{P} - \hat{\mathbf{S}}\hat{\mathbf{S}}^T\|_F^2 + \|\mathbf{S} - \hat{\mathbf{S}}\|_F^2}_{\text{Prior Consistency}}$$

$$\text{s.t. } \mathbf{S} \mathbf{1}_n = \mathbf{1}_m, \mathbf{P} \mathbf{1}_m = \mathbf{1}_m, \text{diag}(\mathbf{P}) = 0, \alpha^T \mathbf{1}_V = 1, \quad (4)$$

where λ and β are regularization parameters, α_v represents the adaptive weight assigned to the v -th view. Specifically, the reconstruction term ensures accurate sample recovery via anchors and the bipartite graph; the manifold regularization term preserves the local topology to refine the learned anchor representation; and the prior consistency term incorporates structural priors to guide the optimization process. By exploiting the synergy among distinct components, the overall clustering performance can be significantly improved.

3 Optimization and Analysis

To address the non-joint convex problem in Eq. (4), an alternating optimization strategy is developed. The detailed procedures for each subproblem are presented as follows:

• **The subproblem w.r.t. $\mathbf{Z}^v$:** By fixing other variables, the optimization problem of Eq. (4) is simplified into:

$$\min_{\mathbf{Z}^v} \sum_{v=1}^V \alpha_v^2 (\|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2 + \lambda \sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij}). \quad (5)$$

Since Eq. (5) is independent across different views, $\mathbf{Z}^v$ can be separately optimized as follows:

$$\min_{\mathbf{Z}^v} \|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2 + \lambda \sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij} \Rightarrow \text{Tr}(\mathbf{S}^T \mathbf{Z}^{vT} \mathbf{Z}^v \mathbf{S} - 2 \mathbf{X}^{vT} \mathbf{Z}^v \mathbf{S}) + 2 \lambda \text{Tr}(\mathbf{Z}^v \mathbf{L}_p \mathbf{Z}^{vT}), \quad (6)$$

where $\mathbf{L}_p$ is the Laplacian matrix derived from $\mathbf{P}$. Setting the derivation of Eq. (6) w.r.t. $\mathbf{Z}^v$ to zero, the optimal solution $\mathbf{Z}^v$ can be calculated as follows:

$$\mathbf{Z}^v = \mathbf{X}^v \mathbf{S}^T (\mathbf{S} \mathbf{S}^T + 2 \lambda \mathbf{L}_p)^{-1}. \quad (7)$$

• **The subproblem w.r.t. $\mathbf{S}$:** By fixing other variables, the optimization problem of Eq. (4) becomes:

$$\min_{\mathbf{S} \mathbf{1}_n = \mathbf{1}_m, \mathbf{S} \geq 0} \sum_{v=1}^V \alpha_v^2 \|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2 + \beta \|\mathbf{S} - \hat{\mathbf{S}}\|_F^2, \quad (8)$$

which can be further transformed into:

$$\min_{\mathbf{S}} \text{Tr}(\mathbf{S}^T (\sum_{v=1}^V \alpha_v^2 \mathbf{Z}^{vT} \mathbf{Z}^v + \beta \mathbf{I}) \mathbf{S}) - 2 \text{Tr}(\mathbf{S}^T (\sum_{v=1}^V \alpha_v^2 \mathbf{Z}^{vT} \mathbf{X}^v + \beta \hat{\mathbf{S}}))$$

$$\text{s.t. } \mathbf{S} \mathbf{1}_n = \mathbf{1}_m, \mathbf{S} \geq 0. \quad (9)$$

Taking the partial derivative w.r.t. $\mathbf{S}$ to zero, the latent solution $\mathbf{S}^*$ can be achieved as follows:

$$\mathbf{S}^* = (\sum_{v=1}^V \alpha_v^2 \mathbf{Z}^{vT} \mathbf{Z}^v + \beta \mathbf{I})^{-1} (\sum_{v=1}^V \alpha_v^2 \mathbf{Z}^{vT} \mathbf{X}^v + \beta \hat{\mathbf{S}}), \quad (10)$$

After that, the optimal solution of $\mathbf{S}$ can be obtained by projecting $\mathbf{S}^*$ onto the feasible region:

$$\min_{\mathbf{S} \mathbf{1}_n = \mathbf{1}_m, \mathbf{S} \geq 0} \|\mathbf{S} - \mathbf{S}^*\|_F^2, \quad (11)$$

Algorithm 1 Optimization procedures for DOTA

Input: Multi-view data $\mathbf{X}^v$, clustering number c and parameters λ and β ;

- 1: Generate the initial anchor matrices $\{\mathbf{Z}^v\}_{v=1}^V$ via K-means on concatenated data;
- 2: Construct the initial bipartite graph $\hat{\mathbf{S}}$ using K-NN strategy [Nie *et al.*, 2014];
- 3: Initialize the view weights $\alpha_v = 1/V$ ($v = 1, \dots, V$);
- 4: **repeat**
- 5: Update $\mathbf{Z}^v$ by Eq. (7);
- 6: Update $\mathbf{S}$ by solving Eq. (11);
- 7: Update $\mathbf{P}$ by solving Eq. (15);
- 8: Update α by solving Eq. (16);
- 9: **until** Eq. (4) converges;

Output: The sample embedding $\mathbf{V}$ derived from bipartite graph, and perform K-means on $\mathbf{V}$ to obtain the clusters.

which can be solved in closed form [Huang *et al.*, 2015].

• **The subproblem w.r.t. $\mathbf{P}$:** By fixing other variables, the optimization problem of Eq. (4) becomes:

$$\begin{aligned} \min_{\mathbf{P}} \lambda \sum_{v=1}^V \alpha_v^2 \sum_{i,j=1}^m \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij} + \beta \|\mathbf{P} - \hat{\mathbf{S}}\hat{\mathbf{S}}^T\|_F^2 \\ \text{s.t. } \mathbf{P}\mathbf{1}_m = \mathbf{1}_m, \mathbf{P} \geq 0, \text{diag}(\mathbf{P}) = 0. \end{aligned} \quad (12)$$

Denoting $\mathbf{D}$ with $d_{ij} = \lambda \sum_{v=1}^V \alpha_v^2 \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2$ and $\mathbf{E} = \hat{\mathbf{S}}\hat{\mathbf{S}}^T$, Eq. (12) can be transformed into:

$$\min_{\mathbf{P} \geq 0, \mathbf{P}\mathbf{1}_m = \mathbf{1}_m, \text{diag}(\mathbf{P}) = 0} \text{Tr}(\mathbf{P}^T \mathbf{D}) + \beta \|\mathbf{P} - \mathbf{E}\|_F^2. \quad (13)$$

Taking the partial derivative w.r.t. $\mathbf{P}$ to zero, the latent solution $\mathbf{P}^*$ can be achieved as follows:

$$\mathbf{P}^* = \mathbf{E} - \frac{1}{2\beta} \mathbf{D}, \quad (14)$$

After that, the optimal solution of $\mathbf{P}$ can be obtained by:

$$\min_{\mathbf{P} \geq 0, \mathbf{P}\mathbf{1}_m = \mathbf{1}_m, \text{diag}(\mathbf{P}) = 0} \|\mathbf{P} - \mathbf{P}^*\|_F^2. \quad (15)$$

• **The subproblem w.r.t. α :** When other variables are fixed, the problem of Eq. (4) is simplified into:

$$\min_{\alpha^T \mathbf{1}_V = 1, \alpha \geq 0} \sum_{v=1}^V \alpha_v^2 h_v, \quad (16)$$

where $h_v = \|\mathbf{X}^v - \mathbf{Z}^v \mathbf{S}\|_F^2 + \lambda \sum_{i,j} \|\mathbf{z}_i^v - \mathbf{z}_j^v\|_2^2 p_{ij}$ denotes the loss of the v -th view. By employing the Lagrange multiplier method, α_v can be obtained as $\alpha_v = \frac{1/h_v}{\sum_{k=1}^V 1/h_k}$.

Efficient Spectral Embedding Transformation

Standard spectral clustering conventionally requires performing eigen-decomposition on the $n \times n$ sample-level similarity graph to derive discriminative embeddings, which severely restricts its scalability for large-scale multi-view tasks. As established in Theorem 1 (the proof is provided in Appendix), DOTA tactfully recovers the global spectral embeddings by propagating spectral information through the learned bipartite graph. Accordingly, DOTA achieves linear scalability w.r.t sample size while effectively preserving the exactness of the underlying global similarity structure.

Dataset	Classes	Data size	Feature size
Leaves	100	1600	192(64/64/64)
COIL100	100	7200	90(30/30/30)
Hdigit	10	10000	1040(784/256)
MNIST	10	20000	459(256/144/59)
KSA	5	20000	279(30/50/30)
CIFAR100	100	50000	3584(512/2048/1024)

Table 2: Detailed information on multi-view datasets.

Theorem 1 (Spectral Recovery Equivalence). *Given the learned bipartite graph $\mathbf{S} \in \mathbb{R}^{m \times n}$, the optimal spectral embedding $\mathbf{V} \in \mathbb{R}^{n \times c}$ of the induced sample-level similarity matrix $\mathbf{S}^T \mathbf{S} \in \mathbb{R}^{n \times n}$ can be exactly recovered by $\mathbf{V} = \mathbf{S}^T \mathbf{U} \mathbf{\Lambda}^{-1/2}$, where $\mathbf{U} \in \mathbb{R}^{m \times c}$ and $\mathbf{\Lambda} \in \mathbb{R}^{c \times c}$ are the top- c eigenvectors and corresponding eigenvalues of the anchor-level matrix $\mathbf{S} \mathbf{S}^T \in \mathbb{R}^{m \times m}$, respectively.*

Computational Complexity

Algorithm 1 summarizes the iterative optimization procedure for solving the objective function in Eq. (4). Specifically, during the initialization stage, generating m anchors and constructing the initial bipartite graph $\hat{\mathbf{S}}$ incur $\mathcal{O}(nm)$. In the optimization phase, updating $\mathbf{Z}^v$ and $\mathbf{S}$ involves computing the inverses of m -order matrices, taking $\mathcal{O}(nmd + nm^2 + m^3)$ and $\mathcal{O}(nm^2 + nmd + m^3)$, respectively. Additionally, the updates for $\mathbf{P}$ and α require $\mathcal{O}(m^2)$ and $\mathcal{O}(nmV^2)$. Finally, the post-processing step involves performing the eigen-decomposition of $\mathbf{S} \mathbf{S}^T$ and executing K-means on the sample embeddings, taking $\mathcal{O}(m^3)$ and $\mathcal{O}(n)$ respectively. Consequently, the overall computational complexity of DOTA is $\mathcal{O}(nm^2 + nmd + m^3)$, which scales linearly with the sample size n .

4 Experiments

4.1 Experimental Settings

To evaluate the effective of the proposed DOTA, comprehensive experiments are conducted on six datasets, including Leaves¹, COIL100 [Zhang *et al.*, 2024c], Hdigit [Jiang *et al.*, 2023], MNIST², KSA and CIFAR100³. The detailed information are summarized in Table 2.

Comparative Methods: DOTA is compared with seven state-of-the-art competitors, including (1) Unified and Discrete Bipartite Graph Learning (**UDBG** [Fang *et al.*, 2023]); (2) Anchor Enhancement strategy with View Correlation (**AEVC** [Liu *et al.*, 2024b]); (3) Cluster-wise Anchor learning for MVC (**CAMVC** [Zhang *et al.*, 2024a]); (4) Robust and Consistent Anchor Graph Learning (**RCAGL** [Liu *et al.*, 2024c]); (5) Anchor Learning with Potential cluster Constraints (**ALPC** [Chen *et al.*, 2025a]); (6) Anchor Graph Learning with Double noise Removal (**AGLDR** [Chen *et al.*, 2025b]); (7) Hypergraph Regularization-based Anchor Learning (**HRAL** [Zeng *et al.*, 2025]).

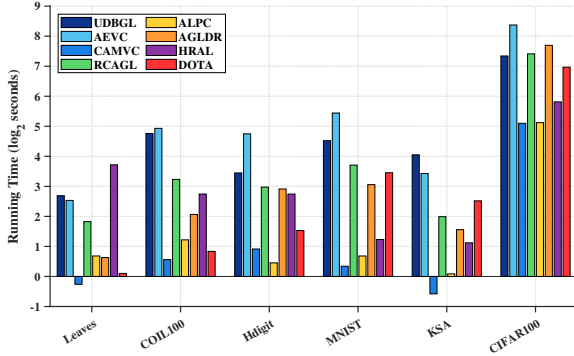
Implementation Details: To ensure fairness, the parameters for all methods are optimized following the suggestions

¹<https://archive.ics.uci.edu/dataset/>

²<https://www.kaggle.com/datasets/hojjatk/mnist-dataset>

³<https://www.cs.toronto.edu/~kriz/cifar.html>

Dataset	UDBGL	AEVC	CAMVC	RCAGL	ALPC	AGLDR	HRAL	DOTA
ACC (%)								
Leaves	65.31 $\pm$ 0.00	79.28 $\pm$ 1.51	<u>81.38</u> $\pm$ 1.84	69.81 $\pm$ 0.00	80.08 $\pm$ 1.93	69.22 $\pm$ 0.92	76.56 $\pm$ 2.05	89.86 $\pm$ 1.21
COIL100	56.47 $\pm$ 0.00	58.59 $\pm$ 1.16	57.39 $\pm$ 1.35	52.65 $\pm$ 0.00	60.22 $\pm$ 1.56	<u>63.21</u> $\pm$ 1.88	49.45 $\pm$ 1.14	79.42 $\pm$ 1.13
Hdigit	84.19 $\pm$ 0.00	85.04 $\pm$ 0.01	87.24 $\pm$ 5.68	78.44 $\pm$ 0.00	78.82 $\pm$ 7.89	<u>87.28</u> $\pm$ 0.02	72.21 $\pm$ 5.62	97.25 $\pm$ 0.00
MNIST	36.57 $\pm$ 0.00	46.35 $\pm$ 0.28	43.30 $\pm$ 1.82	44.02 $\pm$ 0.00	38.46 $\pm$ 2.49	<u>48.13</u> $\pm$ 0.16	44.40 $\pm$ 4.41	52.37 $\pm$ 2.18
KSA	23.70 $\pm$ 0.00	29.66 $\pm$ 0.01	<u>44.32</u> $\pm$ 0.18	21.12 $\pm$ 0.00	30.70 $\pm$ 0.47	23.49 $\pm$ 0.00	34.20 $\pm$ 0.35	45.09 $\pm$ 3.22
CIFAR100	84.55 $\pm$ 0.00	92.96 $\pm$ 1.20	91.04 $\pm$ 1.56	99.97 $\pm$ 0.00	92.22 $\pm$ 1.92	95.04 $\pm$ 0.03	89.57 $\pm$ 1.94	<u>99.92</u> $\pm$ 0.00
Ave. Score	0.5847	0.6531	<u>0.6745</u>	0.6100	0.6342	0.6440	0.6107	0.7732
Ave. Rank	6.83	4.00	4.00	5.67	4.67	<u>3.83</u>	5.83	1.17
NMI (%)								
Leaves	82.98 $\pm$ 0.00	92.06 $\pm$ 0.53	<u>94.04</u> $\pm$ 0.58	87.81 $\pm$ 0.00	92.90 $\pm$ 0.53	86.29 $\pm$ 0.18	91.49 $\pm$ 0.61	96.01 $\pm$ 0.30
COIL100	79.37 $\pm$ 0.00	81.90 $\pm$ 0.42	82.04 $\pm$ 0.37	78.19 $\pm$ 0.00	83.21 $\pm$ 0.49	<u>84.34</u> $\pm$ 0.55	74.90 $\pm$ 0.57	91.55 $\pm$ 0.28
Hdigit	76.38 $\pm$ 0.00	72.25 $\pm$ 0.02	<u>82.20</u> $\pm$ 3.18	73.32 $\pm$ 0.00	78.87 $\pm$ 4.27	75.21 $\pm$ 0.02	66.20 $\pm$ 2.62	93.05 $\pm$ 0.00
MNIST	22.36 $\pm$ 0.00	30.95 $\pm$ 0.10	33.46 $\pm$ 1.15	<u>36.86</u> $\pm$ 0.00	28.70 $\pm$ 0.67	35.65 $\pm$ 0.16	32.25 $\pm$ 2.73	49.40 $\pm$ 0.53
KSA	2.84 $\pm$ 0.00	10.53 $\pm$ 0.02	26.32 $\pm$ 0.60	1.24 $\pm$ 0.00	10.36 $\pm$ 0.24	1.80 $\pm$ 0.00	18.36 $\pm$ 0.01	<u>24.17</u> $\pm$ 1.17
CIFAR100	96.75 $\pm$ 0.00	98.82 $\pm$ 0.21	98.52 $\pm$ 0.25	99.95 $\pm$ 0.00	98.73 $\pm$ 0.30	99.04 $\pm$ 0.02	98.26 $\pm$ 0.32	<u>99.88</u> $\pm$ 0.00
Ave. Score	0.6011	0.6442	<u>0.6943</u>	0.6290	0.6546	0.6372	0.6358	0.7568
Ave. Rank	6.67	5.00	<u>3.17</u>	5.00	4.33	4.50	6.00	1.33
PUR (%)								
Leaves	69.94 $\pm$ 0.00	82.33 $\pm$ 1.16	<u>84.65</u> $\pm$ 1.58	82.94 $\pm$ 0.00	83.13 $\pm$ 1.69	72.60 $\pm$ 0.71	80.03 $\pm$ 1.63	91.03 $\pm$ 1.01
COIL100	60.24 $\pm$ 0.00	62.42 $\pm$ 1.02	61.47 $\pm$ 0.82	<u>67.03</u> $\pm$ 0.00	63.87 $\pm$ 1.37	66.94 $\pm$ 1.39	54.16 $\pm$ 1.59	81.50 $\pm$ 0.85
Hdigit	84.19 $\pm$ 0.00	85.04 $\pm$ 0.01	<u>87.97</u> $\pm$ 4.63	86.97 $\pm$ 0.00	81.69 $\pm$ 6.25	87.28 $\pm$ 0.02	73.48 $\pm$ 4.05	97.25 $\pm$ 0.00
MNIST	39.57 $\pm$ 0.00	49.31 $\pm$ 0.34	47.32 $\pm$ 1.85	57.57 $\pm$ 0.00	42.54 $\pm$ 1.36	51.55 $\pm$ 0.20	47.63 $\pm$ 3.87	<u>57.12</u> $\pm$ 1.22
KSA	23.70 $\pm$ 0.00	29.67 $\pm$ 0.01	<u>44.32</u> $\pm$ 0.18	25.05 $\pm$ 0.00	31.82 $\pm$ 0.01	23.50 $\pm$ 0.00	34.37 $\pm$ 0.05	45.99 $\pm$ 1.98
CIFAR100	87.76 $\pm$ 0.00	94.93 $\pm$ 0.86	93.58 $\pm$ 1.11	99.97 $\pm$ 0.00	94.43 $\pm$ 1.36	95.15 $\pm$ 0.07	92.38 $\pm$ 1.59	<u>99.92</u> $\pm$ 0.00
Ave. Score	0.6090	0.6728	0.6989	<u>0.6992</u>	0.6625	0.6617	0.6368	0.7880
Ave. Rank	7.33	4.67	4.00	<u>3.00</u>	5.00	4.50	6.17	1.33

 Table 3: Clustering performance on six datasets, where the best and second-best results are highlighted in **bold** and underlined, respectively.

 Figure 3: The running time ($\log_2$) comparison across datasets.

in the respective works. For DOTA, the regularization parameters λ and β are tuned via grid search within the range of $\{10^{-3}, 10^{-2}, \dots, 10^3\}$, while the number of anchors is set to $m = 0.1n$. After the optimization procedure, K-means clustering is repeated 20 times on the learned representation to mitigate the impact of random initialization. Finally, the average results on three metrics, including the clustering accuracy (ACC), the normalized mutual information (NMI) and the purity (PUR), are reported to evaluate the performance.

4.2 Comparison Results

The clustering results, including means and standard deviations w.r.t. three metrics, are summarized in Table 3, where the best and second-best values are prominently marked in bold and underlined, respectively. Furthermore, the average scores and ranks of each method across all datasets are calculated to provide a comprehensive comparison. As shown in Table 3, the following conclusions can be drawn: (1) DOTA consistently achieves competitive or superior results across datasets with varying scales and types, where the best average ranks and scores in terms of different metrics are attained, highlighting its effectiveness in multi-view clustering tasks. (2) Compared with two-stage methods (i.e., UDBGL and AEVC), DOTA demonstrates considerable superiority with the average enhancement exceeding 10% across all metrics, thereby confirming that leveraging the interaction between anchor learning and graph construction promotes the overall performance. (3) The results of DOTA are superior to orthogonality-based methods (i.e., RCAGL, ALPC, and AGLDR) and structure-constrained methods (i.e., CAMVC and HRAL), underscoring the effectiveness of simultaneously optimizing the sample-anchor and anchor-anchor similarities.

Meanwhile, to investigate the computational efficiency of DOTA, Fig. 3 displays the running times of all methods on six datasets. Since DOTA and all baselines are anchor-based

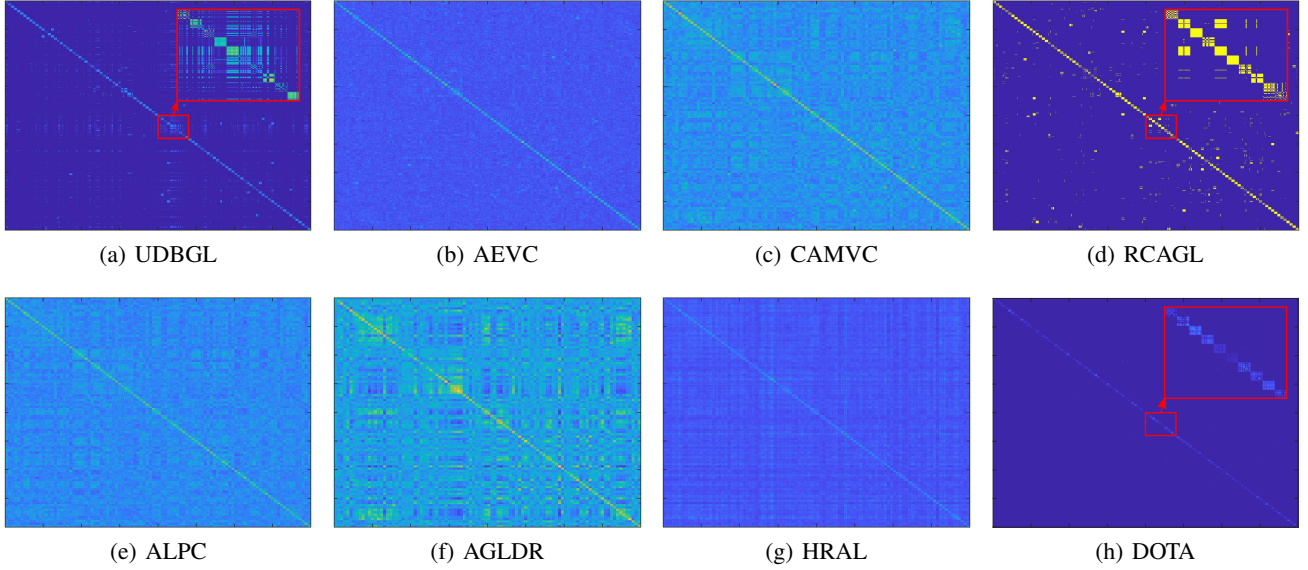


Figure 4: Visual comparison of the learned induced similarity matrices for different methods on Leaves.

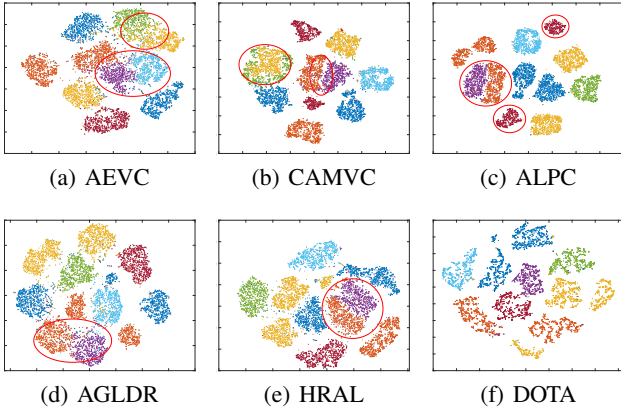


Figure 5: T-SNE results of the latent representation on Hdigit.

MVC methods that have linear time complexity $\mathcal{O}(n)$, they demonstrate strong scalability across all datasets. Empirically, the running time of DOTA surpasses or remains competitive with other methods, demonstrating that our method not only attains excellent clustering results but also preserves highly competitive computational efficiency.

4.3 Visualization Analysis

To intuitively assess the effectiveness of DOTA, the similarity structure and the latent representations are visualized, where the corresponding analysis is presented below:

Similarity Structure: To efficiently capture similarity relationships, bipartite graph learning is widely employed in anchor-based MVC methods. Since pairwise relationships between samples can be characterized by the induced similarity matrix $\mathbf{S}\mathbf{S}^T$, the corresponding visualizations on Leaves are presented in Fig. 4. As illustrated in Figs. 4(a)-(g), the

results of competing methods exhibit numerous irrelevant inter-cluster connections, indicating that the learned bipartite graphs fail to effectively represent the data structure. In contrast, DOTA (Fig. 4h) designs a dual-structure optimization framework to learn the similarity information, so that the intrinsic structure of samples can be captured more accurately.

Latent Representation: To evaluate the discriminative power of the learned embeddings, Fig. 5 illustrates the T-SNE visualizations of the latent representations fed to K-means for various methods on the Hdigit dataset. We can observe that the latent representations of baseline methods (Figs. 5(a)-(e)) exhibit significant overlap between different categories (highlighted by red circles), thereby resulting in ambiguous decision boundaries. Conversely, DOTA (Fig. 5f) reveals a superior distribution, where the intrinsic manifold structure is effectively preserved and prominent margins between different classes are maintained, confirming that the dual-structure optimization framework effectively captures the similarity information and facilitates a more robust partitioning.

4.4 Ablation Study

To evaluate the individual contribution of each component within DOTA, an ablation study is conducted involving three specific variants: DOTA-1 omits the anchor-based manifold regularizer term (i.e., $\lambda = 0$), relying on sample-anchor relation to capture similarity information; DOTA-2 discards the prior consistency term, where the regularization term transforms into $\|\mathbf{P}\|_F^2$ and $\|\mathbf{S}\|_F^2$; and DOTA-3 removes the adaptive view-weighting mechanism by assigning equal importance to all views (i.e., setting $\alpha_v = 1$). The ACC (more results are provided in the Appendix) of DOTA and its simplified versions are presented in Fig. 6, from which the following observations can be drawn: (1) The superiority of DOTA over DOTA-1 validates the effectiveness of dual-structure adaptive optimization, thereby enabling a more discriminative sample

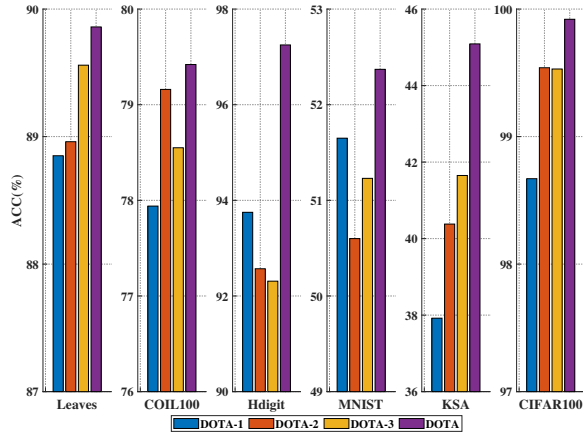


Figure 6: ACC of DOTA and its simplified versions.

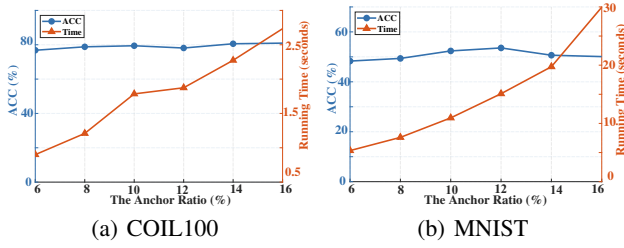


Figure 7: The influence of anchor ratio on ACC and running time.

representation. (2) The observation that DOTA outperforms DOTA-2 underscores the importance of prior consistency in providing effective guidance for the similarity learning process. (3) DOTA achieves better results than DOTA-3, confirming that the adaptive fusion of complementary information across views facilitates the ultimate clustering.

4.5 Parameter Analysis and Convergence

Influence of Anchor Number: To evaluate the sensitivity of DOTA to the anchor scale, Fig. 7 illustrates the trends of ACC and running time across anchor ratios ranging from $\{6\%, 8\%, \dots, 16\%\} \times n$. These observations demonstrate that the computational cost (i.e., orange curve) escalates sharply with the number of anchors, whereas the ACC (i.e., blue curve) remains relatively stable, highlighting that excessive anchors primarily imposes a heavy computational burden without a proportional enhancement in accuracy. Therefore, we empirically set the anchor ratio to $10\%n$ to ensure high efficiency without sacrificing clustering precision.

Parameter Sensitivity: DOTA incorporates two intrinsic regularization parameters to govern the similarity learning process, where λ controls the significance of the anchor-based manifold regularizer and β balances the influence of prior information guidance. Concretely, the parameter sensitivity of DOTA is illustrated in Fig. 8, depicting the ACC and NMI variations across different settings of λ and β . It can be observed that DOTA exhibits high robustness within a broad parameter range and performs better when λ is smaller than 1, underscoring that properly balancing sample-anchor

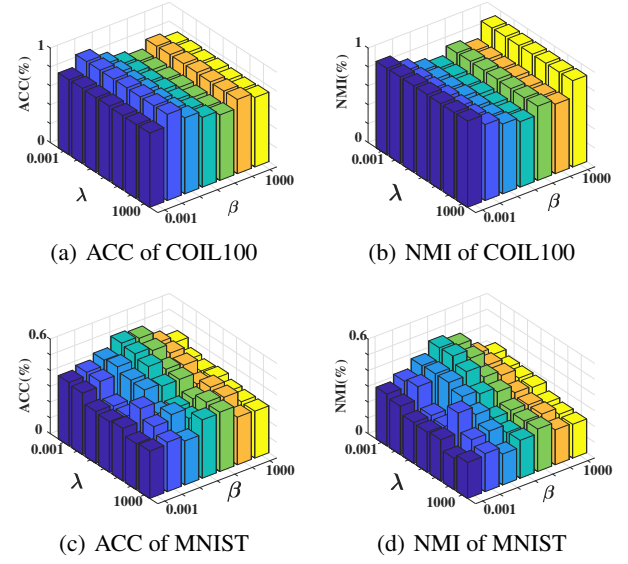
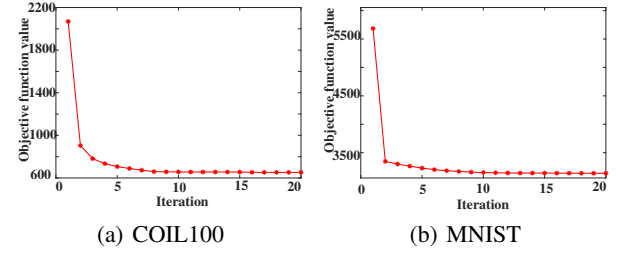

 Figure 8: Clustering performance under varying λ and β .


Figure 9: Variation curves of objective function values.

as well as anchor-anchor structure effectively facilitates the capture of discriminative similarity information.

Convergence Analysis: The convergence of DOTA is shown in Fig. 9, where the objective function values decrease sharply and converge to a steady state within a few iterations, confirming the efficiency of the optimization strategy.

5 Conclusion

In this paper, a novel Dual-tOpology learning with adaptive Anchors (DOTA) method is designed to address the limitations of traditional anchor-based multi-view clustering. Specifically, by leveraging dual-structure optimization of adaptive graphs, DOTA simultaneously characterizes sample-anchor and anchor-anchor relationships, thereby effectively capturing the intrinsic structure and further strengthening the discriminative capacity of embeddings. Meanwhile, DOTA integrates the adaptive view-weights to fuse complementary information and employs an efficient decomposition strategy to learn sample embeddings, which not only achieves robust global structure preservation but also ensures computational efficiency. Extensive experiments underscore the superiority of DOTA when addressing multi-view clustering tasks.

Acknowledgments

This work was supported by the National Natural Science Foundation of China under Grants Nos. 62576161, 72471109, 62276136, 62176116 and 62306282.

References

- [Chen *et al.*, 2025a] Yawei Chen, Huibing Wang, Jinjia Peng, and Yang Wang. Anchor learning with potential cluster constraints for multi-view clustering. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 15939–15947, 2025.
- [Chen *et al.*, 2025b] Zhe Chen, Mingzhi Zhu, Hui Li, and Tianyang Xu. Anchor graph learning with double noise removal for multi-view clustering. *Neural Networks*, page 107779, 2025.
- [Fan *et al.*, 2025] Wentao Fan, Chao Zhang, Chunlin Chen, and Huaxiong Li. Online cross-modal hashing with multi-level memory. In *Proceedings of the ACM International Conference on Multimedia*, pages 6271–6279, 2025.
- [Fan *et al.*, 2026] Wentao Fan, Chao Zhang, Chunlin Chen, and Huaxiong Li. Online cross-modal hashing with expanding label space. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 40, pages 21002–21010, 2026.
- [Fang *et al.*, 2023] Si-Guo Fang, Dong Huang, Xiao-Sha Cai, Chang-Dong Wang, Chaobo He, and Yong Tang. Efficient multi-view clustering via unified and discrete bipartite graph learning. *IEEE Transactions on Neural Networks and Learning Systems*, 35(8):11436–11447, 2023.
- [Huang *et al.*, 2015] Jin Huang, Feiping Nie, and Heng Huang. A new simplex sparse learning model to measure data similarity for clustering. In *Proceedings of the International Joint Conference on Artificial Intelligence*, 2015.
- [Huang *et al.*, 2025] Yanyong Huang, Minghui Lu, Wei Huang, Xiuwen Yi, and Tianrui Li. Time-fs: joint learning of tensorial incomplete multi-view unsupervised feature selection and missing-view imputation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 17503–17510, 2025.
- [Jiang *et al.*, 2023] Bingbing Jiang, Chenglong Zhang, Yan Zhong, Yi Liu, Yingwei Zhang, Xingyu Wu, and Weiguo Sheng. Adaptive collaborative fusion for multi-view semi-supervised classification. *Information Fusion*, 96:37–50, 2023.
- [Jiang *et al.*, 2025a] Bingbing Jiang, Chenglong Zhang, Xinyan Liang, Peng Zhou, Jie Yang, Xingyu Wu, Junyi Guan, Weiping Ding, and Weiguo Sheng. Collaborative similarity fusion and consistency recovery for incomplete multi-view clustering. In *Proceedings of the AAAI conference on artificial intelligence*, volume 39, pages 17617–17625, 2025.
- [Jiang *et al.*, 2025b] Bingbing Jiang, Chenglong Zhang, Zhongli Wang, Xinyan Liang, Peng Zhou, Liang Du, Qinghua Zhang, Weiping Ding, and Yi Liu. Scalable fuzzy clustering with collaborative structure learning and preservation. *IEEE Transactions on Fuzzy Systems*, 33(9):3047–3060, 2025.
- [Jiang *et al.*, 2026] Bingbing Jiang, Jie Wen, Zidong Wang, Weiguo Sheng, Zhiwen Yu, Huanhuan Chen, and Weiping Ding. Scalable semi-supervised learning with discriminative label propagation and correction. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2026.
- [Kou *et al.*, 2025a] Zhiqiang Kou, Si Qin, Hailin Wang, Jing Wang, Mingkun Xie, Shuo Chen, Yuheng Jia, Tongliang Liu, Masashi Sugiyama, and Xin Geng. Label distribution learning with biased annotations assisted by multi-label learning. In *Proceedings of the International Joint Conference on Artificial Intelligence*, pages 16–22, 2025.
- [Kou *et al.*, 2025b] Zhiqiang Kou, Yucheng Xie, Hailin Wang, Jing Wang, Mingkun Xie, Shuo Chen, Yuheng Jia, Tongliang Liu, and Xin Geng. Rankmatch: A novel approach to semi-supervised label distribution learning leveraging rank correlation between labels. In *Proceedings of the Neural Information Processing Systems*, volume 38, pages 158074–158095, 2025.
- [Liang *et al.*, 2025] Xinyan Liang, Shijie Wang, Yuhua Qian, Qian Guo, Liang Du, Bingbing Jiang, Tingjin Luo, and Feijiang Li. Trusted multi-view classification with expert knowledge constraints. In *Proceedings of the International Conference on Machine Learning*, 2025.
- [Liu *et al.*, 2024a] Meng Liu, Yue Liu, Ke Liang, Wenxuan Tu, Siwei Wang, Sihang Zhou, and Xinwang Liu. Deep temporal graph clustering. In *Proceedings of the International Conference on Learning Representations*, 2024.
- [Liu *et al.*, 2024b] Suyuan Liu, Ke Liang, Zhibin Dong, Siwei Wang, Xihong Yang, Sihang Zhou, En Zhu, and Xinwang Liu. Learn from view correlation: An anchor enhancement strategy for multi-view clustering. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 26151–26161, 2024.
- [Liu *et al.*, 2024c] Suyuan Liu, Qing Liao, Siwei Wang, Xinwang Liu, and En Zhu. Robust and consistent anchor graph learning for multi-view clustering. *IEEE Transactions on Knowledge and Data Engineering*, 36(8):4207–4219, 2024.
- [Liu *et al.*, 2025] Meng Liu, Ke Liang, Siwei Wang, Xingchen Hu, Sihang Zhou, and Xinwang Liu. Deep temporal graph clustering: A comprehensive benchmark and datasets. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47(12):11561–11578, 2025.
- [Lu *et al.*, 2025] Zhoumin Lu, Yongbo Yu, Linru Ma, Feiping Nie, and Rong Wang. Capturing individuality and commonality between anchor graphs for multi-view clustering. In *Proceedings of the International Joint Conference on Artificial Intelligence*, pages 5860–5868, 2025.
- [Nie *et al.*, 2014] Feiping Nie, Xiaoqian Wang, and Heng Huang. Clustering and projected clustering with adaptive neighbors. In *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 977–986, 2014.

- [Tang *et al.*, 2025] Chuan Tang, Miaomiao Li, Jun Wang, Chang Tang, Jiahe Jiang, Tianyi Wang, En Zhu, and Xinwang Liu. Multi-view clustering via high-order bipartite graph learning and tensor low-rank representation. *IEEE Transactions on Knowledge and Data Engineering*, 37(11):6521–6534, 2025.
- [Wang *et al.*, 2019] Hao Wang, Yan Yang, and Bing Liu. Gmc: Graph-based multi-view clustering. *IEEE Transactions on Knowledge and Data Engineering*, 32(6):1116–1129, 2019.
- [Wang *et al.*, 2025] Yadi Wang, Xiaoding Guo, Xianhong Hou, Zhijun Miao, Xiaojin Yang, and Jinkai Guo. Multi-modal sentiment recognition with residual gating network and emotion intensity attention. *Neural Networks*, page 107483, 2025.
- [Xu *et al.*, 2024] Deng Xu, Chao Zhang, Zechao Li, Chunlin Chen, and Huaxiong Li. Fast disentangled slim tensor learning for multi-view clustering. *IEEE Transactions on Multimedia*, 27:1254–1265, 2024.
- [Zeng *et al.*, 2025] Yunpeng Zeng, Peng Song, Beihua Yang, Changjia Wang, Guanghao Du, Yanwei Yu, and Wenming Zheng. Hypergraph regularization-based anchor learning for multi-view clustering. *Pattern Recognition*, page 112465, 2025.
- [Zhang *et al.*, 2024a] Chao Zhang, Xiuyi Jia, Zechao Li, Chunlin Chen, and Huaxiong Li. Learning cluster-wise anchors for multi-view clustering. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 16696–16704, 2024.
- [Zhang *et al.*, 2024b] Chenglong Zhang, Yang Fang, Xinyan Liang, Han Zhang, Peng Zhou, Xingyu Wu, Jie Yang, Bingbing Jiang, and Weiguo Sheng. Efficient multi view unsupervised feature selection with adaptive structure learning and inference. In *Proceedings of the International Joint Conference on Artificial Intelligence*, pages 5443–5452, 2024.
- [Zhang *et al.*, 2024c] Chenglong Zhang, Xinjie Zhu, Zidong Wang, Yan Zhong, Weiguo Sheng, Weiping Ding, and Bingbing Jiang. Discriminative multi-view fusion via adaptive regression. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 8(6):3821–3833, 2024.
- [Zhao *et al.*, 2025] Zihua Zhao, Ting Wang, Haonan Xin, Rong Wang, and Feiping Nie. Multi-view clustering via high-order bipartite graph fusion. *Information Fusion*, 113:102630, 2025.