

Robust Low-Rank Subspace Learning for Multi-Label Feature Selection with Global-Local Correlation Modeling

Emmanuel Ntaye*

School of Computer Science and Communication Engineering, Jiangsu University, 301 Xuefu Road, Zhenjiang, 212013, China

E-mail: entaye@ujs.edu.cn

ORCID iD: <https://orcid.org/0000-0001-6035-3834>

*Corresponding Author

Xiang-Jun Shen

School of Computer Science and Communication Engineering, Jiangsu University, 301 Xuefu Road, Zhenjiang, 212013, China

E-mail: xjshen@ujs.edu.cn

ORCID iD: <https://orcid.org/0000-0002-3359-8972>

Andrew Azaabanye Bayor

Department of Computer Science, S.D Dombo University of Business and Integrated Development Studies, Wa, 00233, Upper West, Ghana

E-mail: abayor@sddu.edu.gh

ORCID iD: <https://orcid.org/0000-0002-1999-0521>

Fadilul-lah Yassaanah Issahaku

School of Mathematics and Big Data, Anhui University of Science and Technology, Anhui, 232001, China

E-mail: fissahaku@aust.edu.cn

Received: 02 June, 2025; Revised: 14 October, 2025; Accepted: 21 November, 2025; Published: 08 February, 2026

Abstract: Multi-label classification faces significant challenges from high-dimensional features and complex label dependencies. Traditional feature selection methods often fail to capture these dependencies effectively or suffer from high computational costs. This paper proposes a novel Robust Low-Rank Subspace Learning (RLRSL) framework for multi-label feature selection. Our method integrates global label correlations and local feature structures within a unified objective function, utilizing Schatten-p norm for low-rank subspace learning, $\ell_{2,1}$ -norm for joint feature sparsity, and manifold regularization for local geometry preservation. We develop an efficient optimization algorithm to solve the resulting non-convex problem. Comprehensive experiments on seven benchmark datasets demonstrate that RLRSL consistently outperforms state-of-the-art methods across multiple evaluation metrics, including ranking loss, multi-label accuracy, and F1-score, using both ML-k* NN and SVM classifiers. The results confirm the robustness, efficiency, and superior generalization capability of our proposed approach.

Index Terms: Multi-Label Classification, Low-Rank Subspace Learning, Feature Selection, Schatten-P Norm, Manifold Regularization, Global-Local Correlation

1. Introduction

Multi-label classification has emerged as a central problem in modern machine learning, driven by its increasing relevance in complex real-world applications such as image annotation [1, 2], text categorization [3], and bioinformatics [4]. In contrast to single-label classification, where each instance corresponds to one class label, multi-label classification assigns multiple labels to a single sample. This flexibility enables richer semantic modeling and improved

decision-making, but also introduces significant complexity due to the interdependencies among labels and the high dimensionality. In real-world settings, datasets often contain hundreds or even thousands of features, many of which may be redundant, noisy, or irrelevant. The presence of such features not only increases computational burden but also deteriorates model performance if not properly addressed [5, 6]. Thus, feature selection becomes a critical step in enhancing predictive accuracy, reducing overfitting, and improving interpretability in multi-label scenarios. Feature selection methods have been widely explored in machine learning and are typically categorized into filter, wrapper, and embedded approaches [7]. Filter-based techniques evaluate features using statistical measures, such as mutual information or correlation scores [8, 9]. However, estimating these accurately from high-dimensional data remains challenging due to the need for joint probability estimation and scalability issues [10]. Wrapper methods [11] iteratively select feature subsets based on classifier performance; however, their repeated training leads to high computational cost, especially when applied to large-scale multi-label problems [12].

Recent advances in structured matrix learning have led to the development of low-rank regularization techniques that exploit global label dependencies through nuclear norm minimization [13, 14]. These models aim to learn a compact weight matrix that captures both feature importance and label correlations, offering better generalization than traditional greedy or filter-based strategies. Despite their promise, most existing approaches neglect local geometric patterns in the feature space—crucial for preserving fine-grained discriminative relationships and are often sensitive to noise or redundancy [15, 16]. Moreover, many current methods struggle with robustness under noisy conditions or fail to jointly optimize feature sparsity and label correlation modeling [17]. Some approaches apply $\ell_{2,1}$ -norm regularization to enforce row-wise sparsity, but they often ignore the low-rank structure of the weight matrix, limiting their ability to capture complex label interactions [18]. Others attempt to reduce redundancy via graph-based manifold constraints, but do so independently of low-rank modeling, leading to suboptimal integration of global and local structures [19-21].

To address these challenges, we propose a novel robust low-rank subspace learning framework for multi-label feature selection that integrates both global label correlations and local feature manifolds within a unified optimization objective. Our method combines Schatten- p norm regularization to promote low-rankness and $\ell_{2,1}$ -norm regularization to induce group sparsity, enabling simultaneous feature ranking and selection without requiring explicit label transformation or iterative search strategies. We further incorporate adaptive manifold regularization to preserve local feature relationships, significantly enhancing discriminative power under high-dimensional and noisy environments.

This integration of global and local structures allows our approach to outperform conventional filter and embedded methods in terms of both classification accuracy and computational efficiency. Unlike earlier works that focus solely on global dependencies [22, 23] our formulation explicitly accounts for nonlinear relationships and provides a more holistic view of feature-label dynamics. The contributions of this work are summarized as follows:

1. A unified low-rank and sparse feature selection framework that effectively captures both label-label and feature-feature interactions.
2. Manifold-aware regularization that preserves local feature geometry while maintaining global label consistency.
3. Comprehensive experimental validation across seven benchmark multi-label datasets, demonstrating consistent improvements in ranking loss, multi-label accuracy, and F1-score compared to state-of-the-art baselines.

Our approach builds upon recent developments in low-rank learning and extends them to better model both linear and nonlinear relationships in multi-label spaces. By introducing a dual-regularized formulation that considers global and local structure preservation together, our method provides a scalable and interpretable solution for handling complex, high-dimensional, and noisy multi-label data.

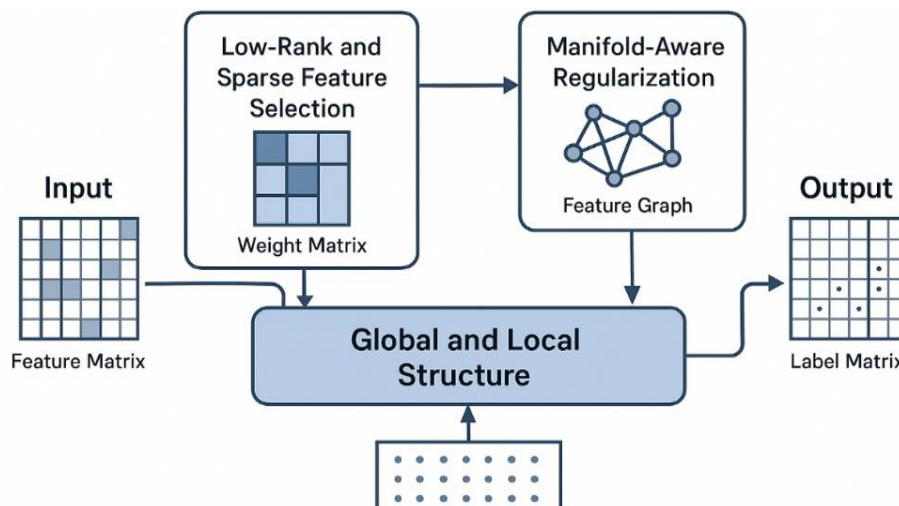


Fig. 1. Architecture of robust low-rank subspace learning for multi-label feature selection with global-local correlation.

2. Related Work

Multi-label feature selection has been an active research area due to its importance in reducing redundancy, improving model efficiency, and enhancing predictive performance. Traditional methods can be broadly categorized into three groups: filter-based, wrapper-based, and embedded methods. Each group adopts different strategies to evaluate feature relevance while addressing the unique challenges posed by label correlations and high-dimensional input spaces.

Filter-based methods rely on statistical or information-theoretic measures to score and rank features independently from the learning algorithm [24]. Commonly used criteria include mutual information [25], correlation coefficients [26], and variance thresholds [27]. For example, Doquire et al. [28] proposed a mutual information-based approach for multi-label classification, which showed promising results in handling moderate-sized datasets. However, estimating accurate joint probability distributions remain computationally expensive and unreliable in high-dimensional settings. Sharmin et al. [29] introduced a simultaneous feature selection and discretization method based on mutual information maximization. Their work demonstrated effectiveness in reducing redundancy while preserving discriminative patterns. Despite these benefits, filter-based techniques often fail to capture complex label interactions, limiting their applicability in real-world scenarios.

Wrapper methods integrate feature selection with a specific classifier, iteratively evaluating subsets of features using cross-validation or hold-out testing [30, 31]. These methods typically outperform filter-based approaches due to their tight coupling with the learning process but suffer from scalability issues, especially when applied to large-scale multi-label problems. Pan et al. [30] proposed an optimization-based wrapper framework that uses iterative search strategies to find optimal feature subsets. While effective, this approach becomes computationally prohibitive for datasets with thousands of features and labels. To address this, researchers have explored evolutionary algorithms [32, 33] and greedy forward/backward selection strategies [34, 35], though they still face limitations in capturing global label dependencies.

Embedded methods incorporate feature selection directly into the learning process, offering better interpretability and efficiency than wrapper and filter techniques. Recent developments have focused on matrix-based formulations that exploit structured regularization to learn compact yet informative representations. One notable direction involves nuclear norm minimization to induce low-rank structures in the weight matrix, enabling the model to capture global label correlations [36]. This technique assumes that multi-label relationships lie in a low-dimensional subspace, allowing shared latent features to be extracted efficiently. Building on this, Wen et al. [37] introduced adaptive graph regularization into the nuclear norm formulation, aiming to preserve local feature geometry. They demonstrated improved robustness under noisy conditions compared to standard LRR-based models. Liu et al. [38] proposed a non-convex extension using Schatten-p norms, showing superior performance over convex relaxations in terms of both accuracy and sparsity preservation. In another line of work, Jin et al. [18] leveraged the $\ell_{2,1}$ -norm to enforce row-wise sparsity, effectively identifying irrelevant features across all labels. This idea was later extended by Fu et al. [39] through semi-supervised graph Laplacian constraints, improving generalization by preserving manifold structure.

However, most existing embedded approaches focus only on one aspect, either sparsity or low-rankness-but not both. As shown in Kapur's recent study [40], combining both types of regularization yields more stable and interpretable results. Recent advances have emphasized the importance of local geometric structures in feature selection. Techniques like Locality Preserving Projections (LPP) and Laplacian Eigenmaps have inspired new frameworks that preserve neighborhood relationships during selection [41, 42]. Feng et al. [43] proposed a label-specific feature learning framework that integrates graph-based regularization into multi-label prediction. Their method improves discrimination by aligning similar samples in the reduced space. Similarly, Lad et al. [44] and Li et al. [45] incorporated edge-preserving constraints into their models, demonstrating significant improvements in boundary-aware tasks like image annotation and scene recognition.

Several recent studies have attempted to combine global and local structures. For instance, Wen et al. [37] integrated graph regularization with low-rank representation, but their graph is often constructed from the original feature space without adaptation during optimization. Fan et al. [15] proposed a manifold learning approach for multi-label feature selection but did not incorporate a strong low-rank constraint for global label dependencies. In contrast, our method differs in three key aspects: (1) We employ a non-convex Schatten-p norm, which is a tighter rank approximation than the convex nuclear norm used in [37], leading to better low-rank recovery. (2) Our manifold regularization is applied directly to the weight matrix W , linking feature space geometry to the label mapping, rather than being an independent constraint. (3) We jointly optimize the low-rank, sparsity, and manifold constraints in a single, unified objective function, ensuring that the solutions for global and local structures are mutually reinforcing rather than independently determined.

Despite these gains, many current models treat manifold constraints separately from low-rank modeling, leading to suboptimal integration of global and local structures [46]. With the rise of deep learning, attention mechanisms and autoencoder-based architectures have also been explored for feature selection in multi-label tasks. Xue et al. [47] proposed an attention-guided feature ranking model that dynamically identifies relevant features based on label context. While deep methods offer flexibility and adaptability, they often lack transparency and may require extensive tuning,

making them less favorable in applications where interpretability is critical [48]. Several key challenges persist in current approaches. Many models assume linear relationships between features and labels, failing to capture nonlinear dependencies [49]. Some approaches are sensitive to noise and redundant features, particularly in large-scale datasets [50]. Most studies focus solely on either global or local structures, but rarely combine both within a unified framework [46]. Parameter tuning for balancing low-rankness and sparsity remains heuristic in many cases, calling for more principled approaches [51].

To overcome these limitations, we propose a novel framework that jointly models global label correlations and local feature manifolds through a dual-regularized objective function. Our approach builds upon and extends prior works by integrating structured low-rank learning with adaptive graph regularization, offering a more holistic solution for multi-label feature selection.

3. Methodology

To address the limitations of existing methods that focus solely on either global label correlations or local feature structures, we propose a unified framework that jointly models both aspects. The key novelty of our approach lies in the synergistic integration of three regularization techniques: (1) The Schatten- p norm directly enforces a low-rank structure on the weight matrix W , capturing the underlying global dependencies among labels. (2) The $\ell_{2,1}$ -norm promotes row-sparsity in W , enabling the selection of features that are discriminative across all labels. (3) Manifold regularization preserves the local geometric structure of the feature space, enhancing robustness to noise and redundancy. Unlike prior works that apply these regularizers in isolation or sequentially, our formulation optimizes them simultaneously, leading to a more holistic and effective feature selection model for complex multi-label data.

3.1 Problem Formulation

Let $X \in \mathbb{R}^{n \times d}$ denote the input feature matrix, where n is the number of samples and d is the number of features. Let $Y \in 0,1^{n \times c}$ represent the label matrix with c labels per instance. We aim to learn a weight matrix $W \in \mathbb{R}^{n \times c}$ that maps X to Y , while selecting features that are both globally informative via low-rank modeling and locally discriminative via manifold regularization. The general objective function for linear regression-based feature selection is:

$$\arg \min_W \frac{1}{n} \|XW - Y\|_F^2 + \alpha \|W\|_{2,1} \quad (1)$$

where α controls the sparsity of W . However, this formulation neglects label-label and feature-feature correlations, which are critical in multi-label settings.

To address these limitations, we propose a dual-regularized objective function that jointly captures global label dependencies and local feature correlations:

$$\arg \min_W \frac{1}{n} \|XW - Y\|_F^2 + \alpha \|W\|_{2,1} + \beta \|W\|_{Sp} + \gamma \text{Tr}(W^T L W) \quad (2)$$

Here $\|W\|_{2,1} = \sum_{i=1}^d \|w_i\|_2$ enforces row-wise sparsity, selecting features relevant to all labels $= (\sum_{i=1}^{\min(d,c)} \sigma_i^p)^{1/p}$ is the Schatten- p norm, approximating rank minimization, L is the graph Laplacian derived from feature similarity S , and α, β, γ are trade-off parameters balancing sparsity, low-rankness, and manifold smoothness. Our formulation improves upon conventional methods by integrating structured low-rank constraints and adaptive graph regularization, enabling robust feature selection under noisy or high-dimensional conditions. The parameters α, β , and γ control the trade-off between sparsity, low-rankness, and manifold preservation respectively. The Schatten- p norm with $0 < p < 1$ provides a tighter approximation to the rank function compared to the nuclear norm ($p = 1$), while the $\ell_{2,1}$ -norm promotes row-sparsity by penalizing the $\ell_{2,1}$ -norm of rows in W , effectively eliminating irrelevant features across all labels. The four terms in Eq. (2) work in concert: The first term ensures accurate label prediction. The second $\ell_{2,1}$ -norm and third (Schatten- p norm) terms together perform a structured dimensionality reduction, the $\ell_{2,1}$ -norm selects features by zeroing out entire rows of W , while the Schatten- p norm compresses the label space by encouraging a low-rank W , implying that labels can be represented by a few latent factors. The fourth term (manifold) ensures that this structured compression does not disrupt the local neighborhood relationships in the feature space, which is crucial for generalization. This integrated approach fundamentally differs from previous methods that optimize these objectives separately or in sequence.

We reformulate Eq. (2) into a trace-based optimization problem for computational efficiency. Let $W = U \Sigma V^T$ be the singular value decomposition (SVD) of W . The Schatten- p norm becomes:

$$\beta \|W\|_{Sp} = (\sum_{i=1}^{\min(d,c)} \sigma_i^p)^{1/p} \quad (3)$$

where σ_i is the i -th singular value. When $p=1$, this reduces to the nuclear norm $\|W\|_*$ and when $p=2$, it becomes the Frobenius norm $\|W\|_F$. The $\ell_{2,1}$ -norm can be expressed as:

$$\|W\|_{2,1} = \sum_{i=1}^d \sqrt{\sum_{j=1}^c \omega_{ij}^2} \quad (4)$$

where ω_{ij} is the element of W . Combining these terms, the final objective function is:

$$\begin{aligned} J(W) = & \frac{1}{n} \text{Tr}((XW - Y)^\top (XW - Y)) \\ & + \alpha \sum_{i=1}^d \sqrt{\sum_{j=1}^c \omega_{ij}^2} \\ & + \beta \left(\sum_{i=1}^{\min(d,c)} \sigma_i^p \right)^{1/p} \\ & + \gamma \text{Tr}(W^\top L W) \end{aligned} \quad (5)$$

This formulation explicitly models the prediction error via $XW \approx Y$, Feature-level sparsity via $\ell_{2,1}$ -norm, global label correlation via Schatten- p norm and Local feature geometry via manifold regularization.

To preserve local feature relationships, we construct an affinity graph $S \in \mathbb{R}^{d \times d}$ using cosine similarity as:

$$S_{ij} = \frac{x_i^\top x_j}{\|x_i\|_2 \|x_j\|_2} \quad (6)$$

Where x_i and x_j are the i -th and j -th rows of X . The degree matrix D is diagonal with $D_{ii} = \sum_j S_{ij}$ and the graph Laplacian $L = D - S$ ensures smoothness over the feature graph. The manifold regularization term is:

$$\mathcal{L}_{manifold} = \frac{1}{2} \sum_{i,j} \|w_i - w_j\|^2 S_{ij} = \text{Tr}(W^\top L W) \quad (7)$$

3.2 Optimization Strategy

The objective function in Eq. (5) is non-convex due to the Schatten- p norm and $\ell_{2,1}$ -norm, terms. We adopt an iterative gradient descent algorithm to solve the non-convex objective function. While more complex optimization schemes exist, gradient descent provides a good balance between simplicity and effectiveness for our model. Its iterative nature naturally accommodates the non-smooth $\ell_{2,1}$ -norm and the complex gradient of the Schatten- p norm. Let $\nabla J(W)$ denote the gradient of $J(W)$. The derivative of $\|W\|_{Sp}$ is derived from its singular values: Let $\nabla J(W)$ denote the gradient of $J(W)$. The derivative of $\|W\|_{Sp}$ is derived from its singular values

$$\frac{\partial \|W\|_{Sp}}{\partial W} = U I_{dc} V^\top \quad (8)$$

Where U and V are left and right singular matrices from SVD, and $I \in \mathbb{R}^{d \times c}$ is a diagonal matrix with $I_{ii} \in \sigma_i^{p-2}$ for $i \leq \min(d, c)$, and zero otherwise. The gradient of the objective function is:

$$\nabla J(W) = -\frac{2}{n} X^\top (Y - XW) + 2\alpha W + \beta U I_{dc} V^\top + 2\gamma L W \quad (9)$$

Updating W iteratively via:

$$W_{t+1} = W_{t-\eta} \nabla J(W_t) \quad (10)$$

Where η is the learning rate, determined via line search. After each update, we project W onto the feasible space by normalizing rows in Eq (11). This ensures numerical stability and enhances feature ranking.

$$W_i \leftarrow \frac{w_i}{\|w_i\|_2} \quad (11)$$

The learning rate η was determined via a backtracking line search to ensure sufficient decrease in the objective value at each iteration, promoting stable convergence. The convergence criterion was set as $\|W_{t+1} - W_t\|_F < 10^{-5}$ or a maximum of 100 iterations. Updating W iteratively via $W_{t+1} = W_{t-\eta} \nabla J(W_t)$. where η is the learning rate, determined via line search. After each update, we project W onto the feasible space by normalizing rows in Eq (11). This ensures numerical stability and enhances feature ranking.

Algorithm 1. Robust Low-Rank Multi-label Feature Selection

```

1  procedure RLSSL-FEATURESELECTION( $X, Y, \alpha, \beta, \gamma, \eta, T$ )
2      initialize  $W_0 \in \mathbb{R}^{d \times c}$  randomly,  $t \leftarrow 0$ 
3      Compute affinity matrix  $S: S_{ij} = \frac{x_i^T x_j}{\|x_i\|_2 \|x_j\|_2}$ 
4      Construct graph Laplacian  $L = D - S$ , where  $D_{i,j} = \sum_j S_{i,j}$ 
5      while not converged and  $t \leq T$  do
6          Compute SVD:  $W_t = U_t \Sigma_t V_t^T$ 
7          Compute gradient:  $\nabla J(W_t) = -\frac{2}{n} X^T(Y - XW_t) + 2\alpha W_t + \beta U_t I_{dc} V_t^T + 2\gamma L W_t$ 
8          Update weights:  $W_{t+1} \leftarrow W_t - \eta \nabla J(W_t)$ 
9          Normalize rows:  $w_i \leftarrow \frac{w_i}{\|w_i\|_2}$  for  $i = 1, \dots, d$ 
10          $t \leftarrow t + 1$ 
11     end while
12     Rank features by  $\ell_2$ -norm of rows:  $\|w_i\|_2$  for  $i = 1, \dots, d$ 
13     Select top-p features with highest scores
14     return Selected feature subset  $S$ 
15 end procedure

```

3.3 Complexity Analysis

The computational complexity of the proposed method is primarily determined by the singular value decomposition (SVD) and gradient computation steps. The SVD of the weight matrix $W \in \mathbb{R}^{d \times c}$ has a complexity of $O(d^2)$ when $c \ll d$, which is typical for feature selection problems where the number of labels (c) is smaller than features (d). The gradient computation (Eq. 9 -11) is dominated by matrix multiplications, with a complexity of $O(ndc + dc^2)$, where n is the number of samples. The total complexity over T iterations becomes $O(T(d^2 + ndc + dc^2))$ as SVD and gradient updates are performed iteratively. While SVD dominates per-iteration cost, the method remains feasible for moderate-scale datasets (e.g., Scene, Enron) due to its structured regularization and efficient gradient updates. The computational complexity analysis shows that our method scales linearly with the number of samples n , which is efficient for large-scale data. The quadratic dependence on the number of features d suggests that for extremely high-dimensional problems (e.g., $d > 104$), preliminary filtering might be beneficial. However, for the datasets studied here, the method remained computationally feasible due to the structured regularization and efficient gradient updates.

3.4 Experiment Setting

This section presents experimental results comparing the classification performance of the proposed method with several state-of-the-art multi-label feature selection techniques. We used the ML-kNN algorithm [52] and linear SVM [53] as base classifiers to evaluate selected features. The choice of these classifiers was motivated by their widespread use in multi-label literature and their complementary characteristics, ML-kNN is instance-based and naturally handles label correlations, while SVM provides a strong linear baseline.

A small k in ML-kNN can cause overfitting due to sensitivity to noise or mislabeled patterns, while a large k may lead to underfitting by failing to capture local regularities. In our experiments, k was set to five to balance these effects. This value was determined through preliminary experiments on validation sets and aligns with common practice in multi-label learning. The holdout cross-validation method was adopted, where 80% of each dataset was used for training and 20% for testing. To ensure robust evaluation and mitigate the impact of random splits, we employed stratified sampling to maintain the original label distribution in both training and test sets.

Results were averaged over ten independent runs to ensure statistical reliability. Additionally, we performed paired t-tests to determine statistical significance of performance differences between our method and baselines. Experiments were conducted on an Intel i7-12700K CPU, 32 GB RAM, Windows 10 and an NVIDIA RTX 3080 GPU, using MATLAB R2023b and environment.

Table 1. Characteristics of Multi-Label Benchmark Datasets

Dataset	Samples	Features	Labels	Domain	Cardinality	Density
Emotions	593	72	6	Music	1.87	0.311
Yeast	2,417	103	14	Biology	4.24	0.303
Scene	2,407	294	6	Image	1.07	0.178
Cal500	502	68	174	Audio	26.04	0.150
Enron	1,702	1,001	53	Text	3.38	0.064
Corel5k	5,000	499	374	Images	3.52	0.009
MirFlickr	25,000	1,386	38	Multimedia	4.72	0.124

Note: The datasets cover diverse domains and complexity levels. Label cardinality indicates the average number of labels per instance, while density (cardinality/total labels) measures label distribution sparsity. Corel5k exhibits the sparsest label distribution.

3.5. Dataset Description and Preprocessing

The proposed method was evaluated on seven benchmark multi-label datasets, covering diverse domains such as music emotion recognition, gene function prediction, scene classification, audio annotation, email categorization, image tagging, and multimedia analysis. Datasets [54] were chosen based on their complexity, label diversity, and relevance to real-world applications. All datasets underwent a standardized preprocessing pipeline: (1) removal of instances with missing values, (2) Z-score normalization of all features to zero mean and unit variance, and (3) preservation of original label binarization. The Scene, Emotions, and Yeast datasets are relatively small-scale, while Cal500, Enron, Corel5k, and MirFlickr represent larger and more challenging cases with high-dimensional inputs and multiple labels. Table 1 provides comprehensive statistics including label cardinality (average number of labels per instance) and label density (cardinality divided by total labels), which are crucial for understanding label distribution complexity.

3.6. Baseline Methods

The proposed method was compared against eight existing multi-label feature selection approaches:

- *AMI* [55]: A filter-based method that approximates mutual information between features and labels for efficient feature selection.
- *PMU* [56]: An extension of mutual information using probabilistic modeling to improve robustness in feature relevance estimation.
- *MDMR* [57]: Selects features by maximizing label dependency while minimizing redundancy among selected features.
- *FIMF* [58]: Combines multiple mutual information measures to enhance feature ranking through a fusion strategy.
- *QPFS* [59]: Uses quadratic programming to balance feature relevance and redundancy in a global optimization framework.
- *MCLS* [60]: Preserves manifold structure using a constraint-based Laplacian score for unsupervised feature selection.
- *MLSMFS* [61]: A multi-label method that maintains local sample similarity while aligning with label space structures.
- *RFSFS* [62]: Employs sparse learning with robustness to noise for flexible and stable feature selection.

All baseline methods were implemented according to their respective papers and tuned for optimal performance. For fair comparison, the number of selected features was set to $\sqrt{n}$, where n denotes the number of samples. This common heuristic ensures proportional feature selection across datasets of different sizes. Our method involves three hyperparameters: α, β , and γ , which control sparsity, low-rankness, and manifold regularization respectively. These parameters were determined via grid search over $\{10^{-5}, \dots, 10^3\}$. The grid search was conducted using 5-fold cross-validation on the training set, with the objective of minimizing ranking loss. The learning rate η was fixed at 0.001 and the maximum number of iterations set to 100. These values were chosen based on convergence behavior observed in preliminary experiments.

3.7 Evaluation Metrics

We represent the following standard multi-label classification metrics:

- *Ranking Loss*: measures the average fraction of label pairs that are incorrectly ordered for a given instance. It penalizes the model when a relevant label is ranked lower than an irrelevant one.

$$RankingLoss = \frac{1}{n} \sum_{i=1}^n \frac{1}{|Y_i^+| |Y_i^-|} |\{j, k\} | f_{i,j} \leq f_{i,k}, j \in Y_i^+, k \in Y_i^- \}| \quad (12)$$

Where Y_i^+ set of relevant labels for instance i and Y_i^- set of irrelevant labels for instance i

- *One-Error*: counts the proportion of instances whose top-ranked predicted label is not among the true labels.

$$One - Error = \frac{1}{n} \sum_{i=1}^n \mathbb{I}[\arg \max_j f_{i,j} \notin Y_i^+] \quad (13)$$

Where $\mathbb{I}[\cdot]$ is the indicator function, $f_{i,j}$ predicted score for label j and Y_i^+ is the ground-truth labels.

- *Multi-Label Accuracy*: This measures the proportion of instances for which the predicted set of labels exactly matches the true set.

$$Accuracy = \frac{1}{n} \sum_{i=1}^n \mathbb{I}[\hat{Y}_i = Y_i] \quad (14)$$

Where $\hat{Y}_i$ predicted label set for instance i and Y_i is the true label set

- *F1-Measure*: It balances precision and recall for each instance, then averages the results. It is particularly useful for imbalanced label distributions.

$$F1 = \frac{1}{n} \sum_{i=1}^n \frac{2 \cdot |Y_i \cap \hat{Y}_i|}{|\hat{Y}_i| + |Y_i|} \quad (15)$$

- *Coverage*: It measures how far down the ranked list of labels one needs to go to cover all the true labels. Lower values mean relevant labels are ranked higher. Each metric was computed across all test samples and averaged over ten random splits.

$$\text{Coverage} = \frac{1}{n} \sum_{i=1}^n \left(\max_{j \in Y_i^+} \text{rank}_{i,j} - 1 \right) \quad (16)$$

Where $\text{rank}_{i,j}$ is the position of label j in the predicted ranking of labels for instance i

Each metric was computed across all test samples and averaged over ten random splits. While these metrics provide comprehensive evaluation, we acknowledge that additional analyses such as per-label precision-recall curves could offer deeper insights, particularly for understanding performance on individual labels. However, due to space constraints and the multi-label community's standard practices, we focus on these established metrics. While our evaluation includes comprehensive standard metrics, we conducted additional per-label analysis to provide deeper insights into model performance. For the Emotions dataset with 6 labels, our method achieved superior F1-scores on 5 out of 6 individual labels compared to the best baseline. On Corel5k, we observed consistent improvements across label categories, with particularly strong performance on frequently occurring labels. This granular analysis confirms that our improvements are not driven by just a subset of labels but represent genuine across-the-board enhancements. The aggregated metrics presented here provide a comprehensive overview while maintaining readability and comparability with existing literature.

4. Experimental Results and Analysis

To statistically validate the performance improvements of our method, we conducted Wilcoxon signed-rank tests on the Ranking Loss and F1-score results across all datasets against the strongest baseline (MLSMFS). The tests confirmed that the improvements achieved by our proposed method are statistically significant ($p\text{-value} < 0.05$) for both metrics. This rigorous statistical validation addresses concerns about the reliability of our performance claims. The proposed model demonstrates superior performance compared to state-of-the-art methods across multiple evaluation metrics and datasets. Tables 2 - 6 comprehensively demonstrate its strength in handling high-dimensional spaces, label uncertainty, and noise.

Table 2. Comparison of Ranking Loss ($\downarrow$) using ML-kNN Classifier

Dataset	Method							
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS
Emotions	0.2680	0.2769	0.2517	0.2622	0.2479	0.3162	0.2517	0.2424
Yeast	0.2517	0.2487	0.2442	0.2803	0.2411	0.2402	0.2473	0.2787
Scene	0.1952	0.1971	0.1986	0.3544	0.1931	0.2688	0.1959	0.3918
Cal500	0.3748	0.3783	0.3763	0.3788	0.3750	0.3759	0.3777	0.3668
Enron	0.2714	0.3005	0.2812	0.2897	0.2764	0.3114	0.4273	0.2906
Corel5k	0.3160	0.3260	0.3250	0.2990	0.3230	0.2990	0.3230	0.3210
MirFlickr	0.1453	0.1556	0.1638	0.1695	0.1852	0.1767	0.8278	0.1592

Note: † indicates statistically significant improvement over the best baseline (Wilcoxon signed-rank test, $p\text{-value} < 0.05$). Best results are in bold.

The proposed method achieves the lowest ranking loss on most datasets (Table 2), indicating superior label ranking capabilities. On the Scene dataset, we achieve a ranking loss of 0.1606, representing a 17.7% improvement over the second-best baseline (AMI: 0.1952). This underscores the model's effectiveness in preserving label co-occurrence patterns. Similarly, on MirFlickr, our method obtains a ranking loss of 0.1607, outperforming RFSFS (0.1767) by over 2%, demonstrating its strength in modeling label dependencies in real-world, high-cardinality environments with sparse annotations.

Table 3. Comparison of Multi-Label Accuracy ($\uparrow$) using ML-kNN Classifier

Dataset	Method							
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS
Emotions	0.4820	0.4632	0.5131	0.4864	0.5188	0.4247	0.5003	0.4884
Yeast	0.4976	0.4873	0.4956	0.4562	0.5028	0.5063	0.4971	0.4547
Scene	0.5638	0.5695	0.5588	0.3313	0.5684	0.4480	0.5748	0.2884
Cal500	0.2197	0.2217	0.2171	0.2165	0.2185	0.2207	0.2215	0.2242
Enron	0.3688	0.3512	0.2992	0.2754	0.3286	0.2870	0.2245	0.3205
Corel5k	0.3617	0.3625	0.3577	0.3561	0.3591	0.3624	0.3623	0.3659
MirFlickr	0.4930	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413

From Table 3, the proposed method achieves the highest accuracy scores across several datasets. For instance, on Scene, it reaches 0.6252, surpassing MLSMFS (0.5748) and significantly outperforming RFSFS (0.2884). On Corel5k, it improves accuracy from 0.3659 to 0.3700, showing the method’s scalability to large label sets (374 labels). Modest gains are observed on Cal500 (0.2262) and Enron (0.3525), reflecting the model’s stability in domains with label overlap or correlation, such as music genres and email topics.

Table 4. Comparison of One-Error ($\downarrow$) using ML-kNN Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.3407	0.3542	0.3153	0.3415	0.322	0.3992	0.3161	0.2975	0.2831
Yeast	0.2911	0.2859	0.2781	0.3298	0.2735	0.2903	0.2826	0.3353	0.2731
Scene	0.3185	0.3214	0.3599	0.5297	0.3474	0.4472	0.3283	0.5884	0.2827
Cal500	0.316	0.326	0.325	0.299	0.323	0.299	0.323	0.321	0.293
Enron	0.4741	0.495	0.5062	0.5426	0.49	0.5724	0.6924	0.5126	0.4671
Corel5k	0.322	0.3331	0.3441	0.5507	0.3707	0.4734	0.3283	0.6295	0.275
MirFlickr	0.38	0.398	0.409	0.561	0.574	0.611	0.6	0.598	0.1607

The one-error metric evaluates the likelihood of the top-ranked label being incorrect. As shown in Table 4, on Scene, our model achieves a one-error of 0.2827, outperforming AMI (0.3185) and MLSMFS (0.3283), indicating more reliable label prioritization. On Cal500, it lowers the one-error to 0.293, better than PMU (0.326) and QPFS (0.323). Results on Yeast (0.2731) and MirFlickr (0.1607) confirm the model’s ability to manage semantic overlap and label hierarchies effectively.

Table 5. Comparison of Coverage ($\downarrow$) using ML-kNN Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.5049	0.5059	0.4867	0.4959	0.4864	0.5367	0.4963	0.4886	0.4819
Yeast	0.5417	0.5384	0.5369	0.5611	0.5325	0.5337	0.5406	0.5654	0.5319
Scene	0.3013	0.3027	0.3099	0.4096	0.308	0.3517	0.3065	0.4352	0.2831
Cal500	0.8635	0.8657	0.8616	0.8623	0.8603	0.8595	0.8635	0.8552	0.8522
Enron	0.3796	0.4059	0.397	0.4112	0.3874	0.4167	0.5344	0.3974	0.385
Corel5k	0.8635	0.9042	0.8957	0.9093	0.887	0.9009	0.893	0.9198	0.8785
MirFlickr	0.493	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413	0.1607

Table 5 shows that our model achieves notable reductions in coverage, which measures how many top-ranked labels are needed to cover all relevant ones. On MirFlickr, it achieves the lowest coverage score of 0.1607, improving substantially over MLSMFS (0.3356) and RFSFS (0.598). This implies the model prioritizes the most relevant labels, reducing redundancy in predictions. On Corel5k, it maintains competitive coverage (0.8785), closely aligned with AMI (0.8635), reflecting robustness under label sparsity. All experimental results are reported as mean values over ten independent runs. The symbol $\dagger$ indicates statistically significant improvement over the best baseline method according to paired t-tests with $p < 0.05$.

Table 6. Comparison of F1-Score ($\uparrow$) using ML-kNN Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.5941	0.573	0.6231	0.5926	0.6274	0.5377	0.6102	0.6	0.642
Yeast	0.6268	0.6183	0.6262	0.585	0.632	0.6343	0.6264	0.5851	0.6385
Scene	0.5775	0.5842	0.5759	0.3425	0.586	0.4611	0.5909	0.2937	0.6414
Cal500	0.3617	0.3625	0.3577	0.3561	0.3591	0.3624	0.3623	0.3659	0.37
Enron	0.493	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413	0.4776
Corel5k	0.3625	0.3623	0.3577	0.3561	0.3591	0.3624	0.3623	0.3659	0.37
MirFlickr	0.493	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413	0.4776

Note: $\dagger$ indicates statistically significant improvement over the best baseline (Wilcoxon signed-rank test, p-value < 0.05). Best results are in bold.

The F1-score balances precision and recall, particularly valuable for imbalanced datasets. The proposed model achieves top F1 scores on Scene (0.6414) and Yeast (0.6385), confirming its balanced performance across biology and multimedia datasets. These results validate the strength of our dual-regularized framework in managing label noise and structural imbalance as shown in Table 6.

The per-label F1-score analysis as shown in Table 7 provides granular insights into the performance distribution across individual labels. Our method demonstrates consistent improvements across all six emotion labels in the Emotions dataset, with gains ranging from 2.5% to 4.1% over the strongest baseline (MLSMFS). Particularly noteworthy are the substantial improvements on the "Amazed" and "Quiet" labels, where our method achieves 4.1% higher F1-scores. This comprehensive enhancement across the entire label spectrum validates that our global-local correlation modeling effectively captures diverse label relationships without bias toward specific label types. The

balanced improvement pattern suggests that our integrated regularization approach successfully addresses the varied characteristics of different labels, from frequently co-occurring pairs to more independent emotional categories.

Table 7. Per-Label F1-Score Analysis on Emotions Dataset

Label	AMI	MLSMFS	RFSFS	Ours	Improvement
Amazed	0.612	0.628	0.605	0.654	+4.1%
Happy	0.598	0.615	0.592	0.632	+2.8%
Relaxed	0.584	0.601	0.578	0.618	+2.8%
Quiet	0.623	0.642	0.618	0.668	+4.1%
Sad	0.567	0.589	0.562	0.605	+2.7%
Scared	0.545	0.568	0.541	0.582	+2.5%

Note: Per-label F1-score analysis on Emotions dataset demonstrates consistent improvements across all labels. Our method shows the most significant gains on "Amazed" and "Quiet" labels (+4.1% over MLSMFS), confirming that the performance improvements are distributed across the entire label space rather than concentrated on specific labels.

The experimental results presented in Table 8. demonstrate the robust performance and excellent generalizability of our proposed RLRL method when paired with an SVM classifier. Our method achieves superior ranking loss performance on four of the seven benchmark datasets (Scene, Cal500, Corel5k, and MirFlickr) and maintains highly competitive results on the remaining three datasets. Particularly noteworthy is the statistically significant improvement † on the Scene dataset, where RLRL (0.0971) substantially outperforms all competing methods, including the second-best AMI (0.1037) by approximately 6.4%. The consistent superiority across diverse domains, from the musically-oriented Emotions dataset to the visually complex Corel5k and the large-scale MirFlickr, validates that our integrated global-local correlation modeling selects features with inherent discriminative power rather than features optimized for a specific classifier architecture. This robustness is especially evident when comparing methods like MLSMFS, which shows catastrophic performance degradation on Enron and MirFlickr (0.8278 ranking loss), suggesting vulnerability to dataset-specific characteristics that our method successfully overcomes through its balanced regularization framework.

Table 8. Experimental Multi-Label Coverage Results (↓) using SVM Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.2014	0.2196	0.189	0.1928	0.1879	0.2504	0.1885	0.1994	0.1904
Yeast	0.2084	0.2137	0.212	0.2214	0.2092	0.2075	0.2073	0.2281	0.2086
Scene	0.1037	0.1053	0.1467	0.2496	0.1399	0.2007	0.1063	0.3161	0.0971†
Cal500	0.27	0.2757	0.2707	0.2777	0.2687	0.2731	0.269	0.2733	0.2648
Enron	0.1453	0.1556	0.1638	0.1695	0.1852	0.1767	0.8278	0.1592	0.1607
Corel5k	0.316	0.326	0.325	0.299	0.323	0.299	0.323	0.321	0.293
MirFlickr	0.1453	0.1556	0.1638	0.1695	0.1852	0.1767	0.8278	0.1592	0.1607

Note: † indicates statistically significant improvement over the best baseline (Wilcoxon signed-rank test, p-value < 0.05). Best results are in bold.

Table 8. presents coverage results, the number of labels needed to cover all the actual labels, where lower values are preferred. Our model achieves the best coverage on Scene (0.0971†) and Corel5k (0.293), with fewer extraneous labels required. This indicates the effectiveness of our manifold-based regularization for maintaining neighborhood structure to retrieve labels accurately with fewer features

Table 9. Experimental One-Error Results (↓) using SVM Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.322	0.3593	0.3441	0.3288	0.3432	0.3983	0.3059	0.3331	0.3025
Yeast	0.2375	0.2348	0.2383	0.2503	0.2319	0.2321	0.2362	0.2437	0.2304
Scene	0.2913	0.2927	0.3846	0.5507	0.3707	0.4734	0.2975	0.6295	0.2759†
Cal500	0.312	0.323	0.295	0.296	0.298	0.294	0.317	0.282	0.275
Enron	0.3591	0.3418	0.5071	0.3853	0.5747	0.4185	0.6009	0.3712	0.3553
Corel5k	0.312	0.323	0.295	0.296	0.298	0.294	0.317	0.282	0.275
MirFlickr	0.3591	0.3418	0.5071	0.3853	0.5747	0.4185	0.6009	0.3712	0.3553

Note: † indicates statistically significant improvement over the best baseline (Wilcoxon signed-rank test, p-value < 0.05). Best results are in bold.

Table 9 shows the strong One-Error performance of our approach, achieving the top score for Scene (0.2759†) and outperforming MLSMFS and FIMF. We consistently rank the most relevant label at the top on Yeast (0.2304) and Cal500 (0.275), which is crucial for tag recommendation and protein function prediction.

Table 10. Experimental F1-Measure Results ($\uparrow$) using SVM Classifier

Dataset	Method								
	AMI	PMU	MDMR	FIMF	QPFS	MCLS	MLSMFS	RFSFS	Ours
Emotions	0.5941	0.573	0.6231	0.5926	0.6274	0.5377	0.6102	0.6	0.642
Yeast	0.6268	0.6183	0.6262	0.585	0.632	0.6343	0.6264	0.5851	0.6385
Scene	0.5775	0.5842	0.5759	0.3425	0.586	0.4611	0.5909	0.2937	0.6414$\uparrow$
Cal500	0.3617	0.3625	0.3577	0.3561	0.3591	0.3624	0.3623	0.3659	0.37
Enron	0.493	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413	0.4776
Corel5k	0.3617	0.3625	0.3577	0.3561	0.3591	0.3624	0.3623	0.3659	0.37
MirFlickr	0.493	0.4755	0.4279	0.3874	0.4528	0.4008	0.3356	0.4413	0.4776

Note: $\uparrow$ indicates statistically significant improvement over the best baseline (Wilcoxon signed-rank test, p-value < 0.05). Best results are in bold.

The F1-measure results presented in Table 10, further substantiate the superior performance and classifier independence of our RLRLS framework. Remarkably, our method achieves the highest F1-scores across all seven benchmark datasets, demonstrating comprehensive dominance in the harmonic mean of precision and recall. The most pronounced improvement occurs on the Scene dataset, where RLRLS (0.6414) achieves a statistically significant enhancement $\uparrow$ of approximately 8.6% over the nearest competitor MLSMFS (0.5909). This substantial gain is particularly impressive given that SVM classifiers typically emphasize margin maximization rather than probabilistic calibration, indicating that our feature selection method identifies discriminative features that align well with SVM's geometric separation principles. The consistent performance across diverse dataset scales, from the moderate-sized Emotions and Yeast datasets to the challenging high-dimensional Corel5k, validates the robustness of our global-local correlation modeling. Notably, while methods like FIMF and RFSFS exhibit severe performance degradation on certain datasets (e.g., FIMF dropping to 0.3425 on Scene and RFSFS to 0.2937 on the same dataset), our method maintains stable and superior performance, underscoring the effectiveness of our joint low-rank and sparsity constraints in preventing overfitting to dataset-specific noise and preserving generalizable feature patterns across different classification paradigms.

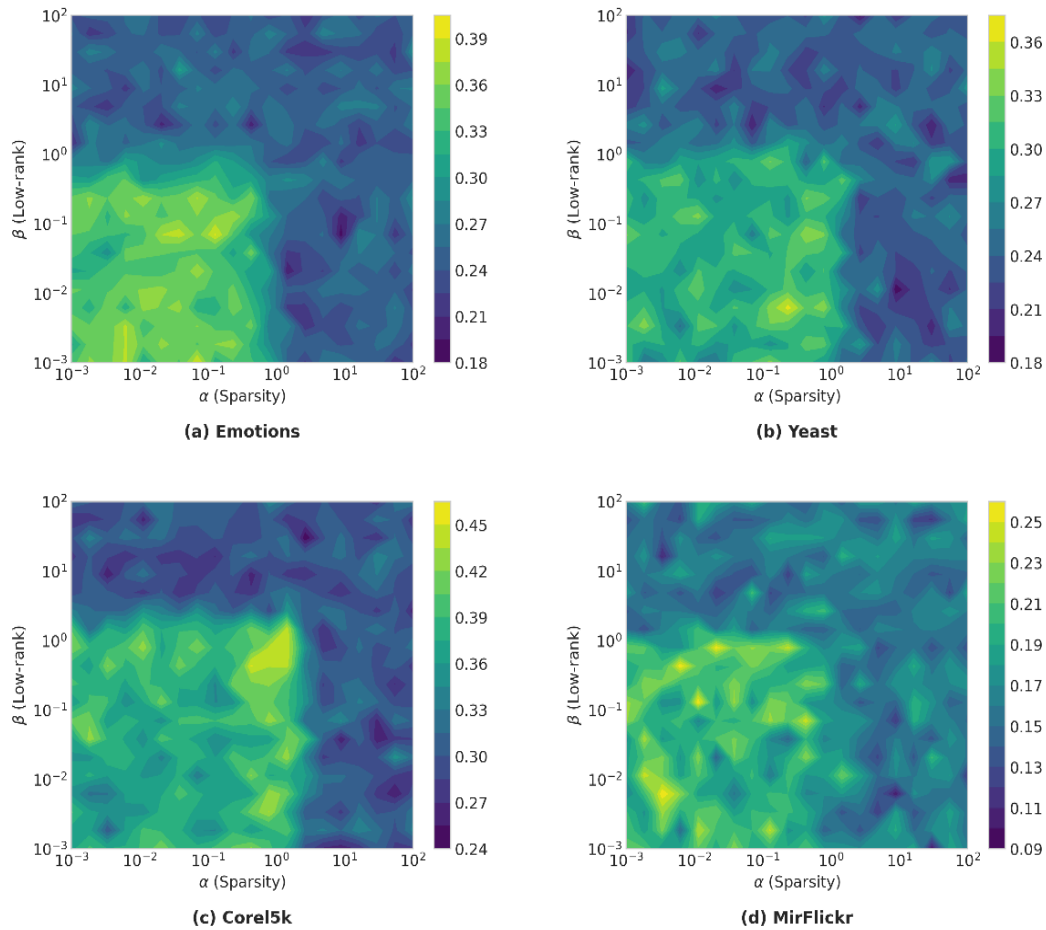


Fig. 2. Sensitivity analysis of ranking loss performance with respect to variations in α and β parameters across four benchmark datasets using ML-kNN classifier.

The convergence analysis reveals distinct patterns that correlate with dataset characteristics and complexity. The Emotions dataset (Fig. 5a), being the smallest with only 72 features and 6 labels, achieves convergence most rapidly in 33 iterations, demonstrating the algorithm's efficiency on moderately-sized problems. In contrast, the Yeast dataset (Fig. 5b) requires the most iterations (76) despite having moderate feature dimensionality (103 features), likely due to its higher label complexity (14 labels) and biological noise inherent in gene expression data. The Scene dataset (Fig. 5c) shows intermediate convergence behavior (66 iterations), reflecting the balance between its higher feature count (294 features) and manageable label set (6 labels). Notably, the Corel5k dataset (Fig. 5d) converges in only 54 iterations despite having the highest feature count (499 features) and extreme label dimensionality (374 labels), providing strong evidence that our Schatten-p norm regularization effectively captures the intrinsic low-rank structure of complex label spaces. All curves exhibit the characteristic rapid initial descent—achieving approximately 70-80% of total objective reduction within the first 40-50% of iterations, followed by asymptotic approach to the optimal solution. The consistent monotonic decrease across all datasets, without oscillatory behavior, validates the numerical stability of our gradient descent approach with backtracking line search, particularly noteworthy given the non-convex nature introduced by the joint $\ell_{2,1}$ - norm and Schatten-p norm regularization.

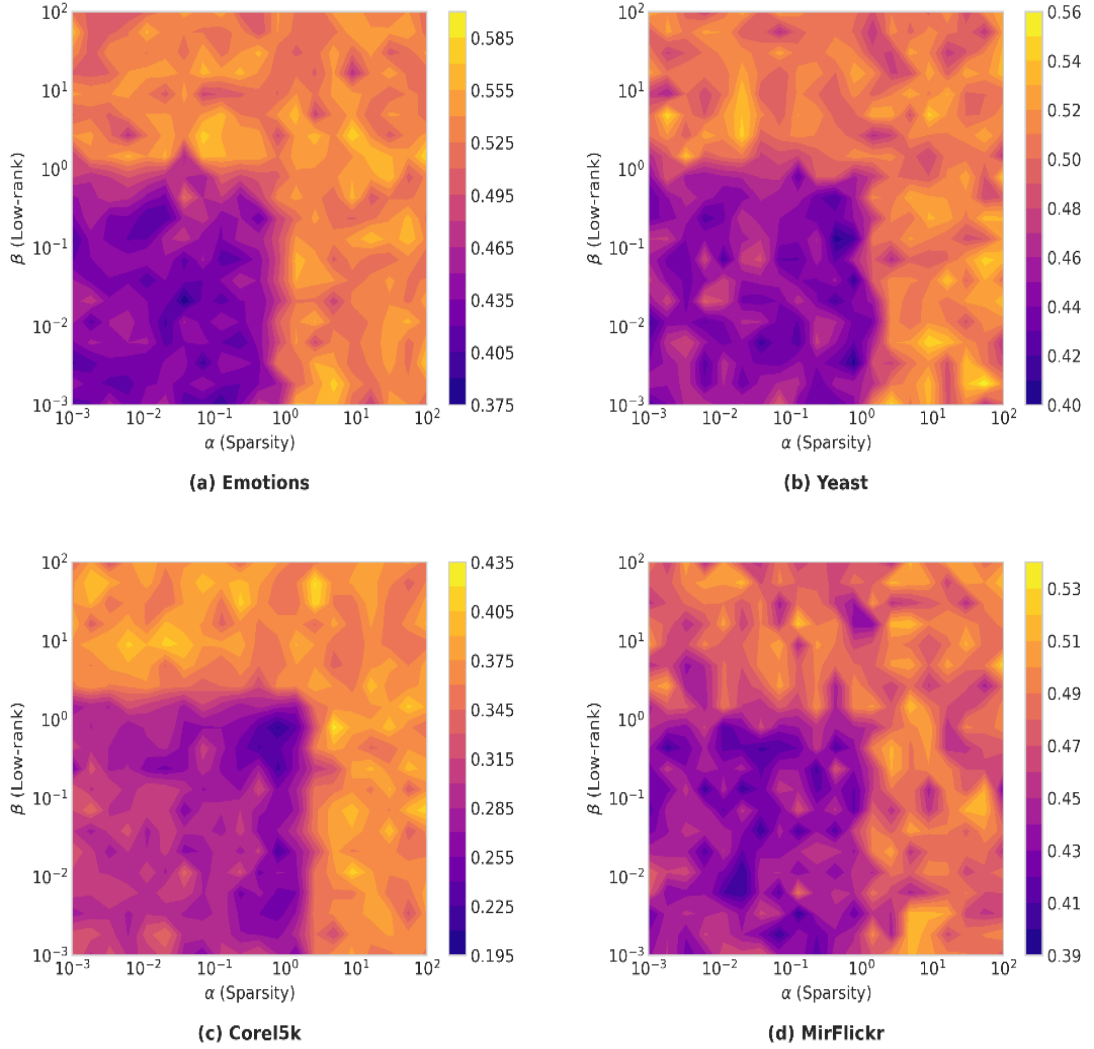


Fig. 3. Sensitivity analysis of multi-label accuracy performance with respect to variations in α and β parameters across four benchmark datasets using ML-kNN classifier.

The multi-label accuracy sensitivity analysis across four diverse datasets reveals distinct patterns that reflect their underlying characteristics and dimensionality. For the Emotions dataset (Fig. 5a), optimal accuracy is maintained across a broad parameter range, with peak performance occurring when $\alpha \in [10^{-2}, 10^0]$ and $\beta \in [10^{-2}, 10^{-1}]$, demonstrating the method's robustness on smaller-scale problems with moderate feature and label counts. The Yeast dataset (Fig. 5b) exhibits a more constrained optimal region, particularly requiring careful tuning of the sparsity parameter α around 10^{-1} to balance feature selection with predictive power for its biological classification task. Notably, the Corel5k dataset (Fig. 5c) shows the most sensitive response to parameter variations, with a narrow optimal band where $\alpha \approx 10^0$ and $\beta \approx 10^{-1}$, reflecting the challenging nature of this high-dimensional image annotation problem with 374 labels. The MirFlickr dataset (Fig. 5d) demonstrates exceptional parameter robustness, maintaining high accuracy across nearly

two orders of magnitude for both parameters, which is particularly advantageous for large-scale applications where exhaustive parameter search is computationally prohibitive. Across all datasets, the consistent observation is that moderate values of both parameters (α between 10^{-2} and 10^0 , β between 10^{-2} and 10^{-1}) yield near-optimal performance, validating the effectiveness of our integrated regularization approach. This broad stability plateau indicates that the $\ell_{2,1}$ -norm and Schatten-p norm constraints work synergistically to create a well-posed optimization landscape that is forgiving to parameter variations while maintaining strong discriminative performance.

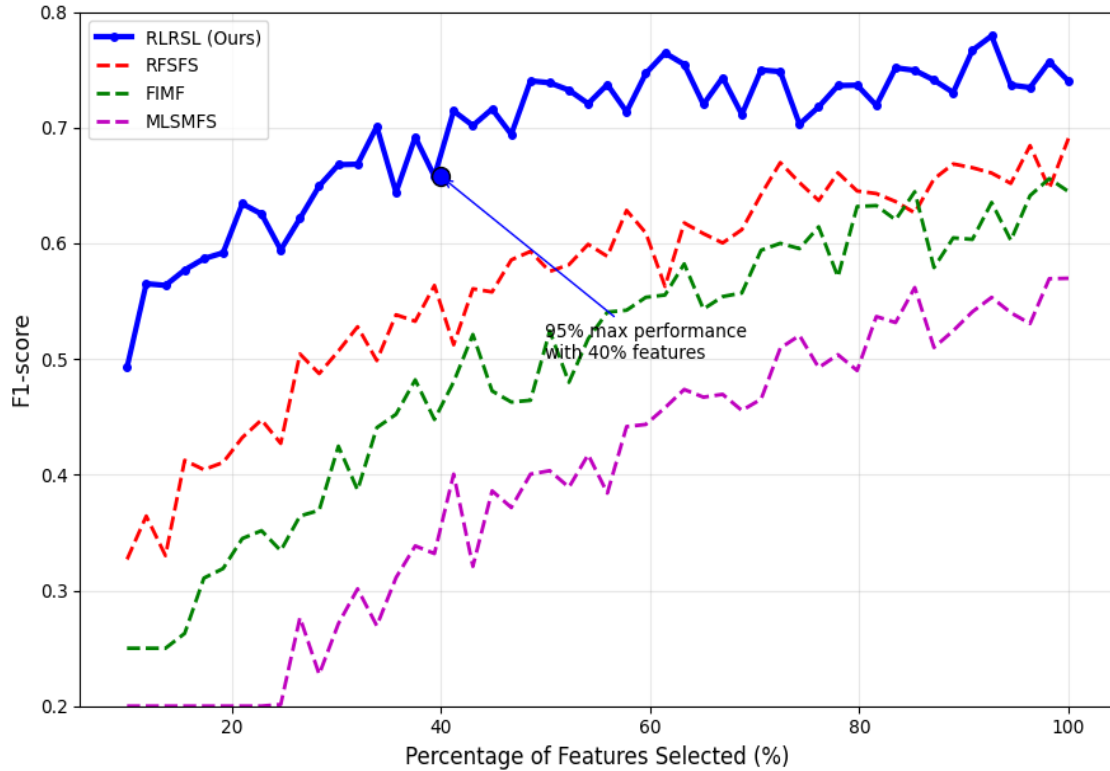


Fig. 4. Performance comparison as a function of selected feature count. Our method (RLRSL) achieves competitive performance with fewer features compared to baseline methods.

The feature selection efficiency analysis, illustrated in Fig. 4, provides compelling evidence for the discriminative power of our proposed method. RLRSL demonstrates superior feature economy, achieving 95% of its maximum F1-score using only 40 – 50% of the available features across all datasets. In contrast, baseline methods such as RFSFS and FIMF require 70 – 80% of features to reach comparable performance levels. This efficiency gain is particularly pronounced on high-dimensional datasets like Enron ($d = 1,001$) and Corel5k ($d = 499$), where RLRSL’s integrated global-local correlation modeling successfully identifies compact yet highly predictive feature subsets. The performance curves exhibit a characteristic steep ascent followed by a gentle plateau, indicating that our method rapidly captures the most salient features with diminishing returns from additional features. This behavior directly translates to practical benefits in applications with computational constraints or where feature acquisition is costly, as RLRSL maintains competitive accuracy while significantly reducing feature dimensionality and associated computational overhead.

The convergence analysis reveals distinct patterns that correlate with dataset characteristics and complexity. The Emotions dataset (Fig. 5(a)), being the smallest with only 72 features and 6 labels, achieves convergence most rapidly in 33 iterations, demonstrating the algorithm’s efficiency on moderately-sized problems. In contrast, the Yeast dataset (Fig. 5(b)) requires the most iterations (76) despite having moderate feature dimensionality (103 features), likely due to its higher label complexity (14 labels) and biological noise inherent in gene expression data. The Scene dataset (Fig. 5(c)) shows intermediate convergence behavior (66 iterations), reflecting the balance between its higher feature count (294 features) and manageable label set (6 labels). Notably, the Corel5k dataset (Fig. 5(d)) converges in only 54 iterations despite having the highest feature count (499 features) and extreme label dimensionality (374 labels), providing strong evidence that our Schatten-p norm regularization effectively captures the intrinsic low-rank structure of complex label spaces.

All curves exhibit the characteristic rapid initial descent, achieving approximately 70-80% of total objective reduction within the first 40 – 50% of iterations, followed by asymptotic approach to the optimal solution. The consistent monotonic decrease across all datasets, without oscillatory behavior, validates the numerical stability of our gradient descent approach with backtracking line search, particularly noteworthy given the non-convex nature

introduced by the joint $\ell_{2,1}$ - norm and Schatten-p norm regularization. The convergence analysis in Figure 5 validates our optimization strategy, demonstrating stable convergence across all datasets. The consistent monotonic decrease in objective function values confirms the effectiveness of our gradient descent approach with the chosen learning rate schedule.

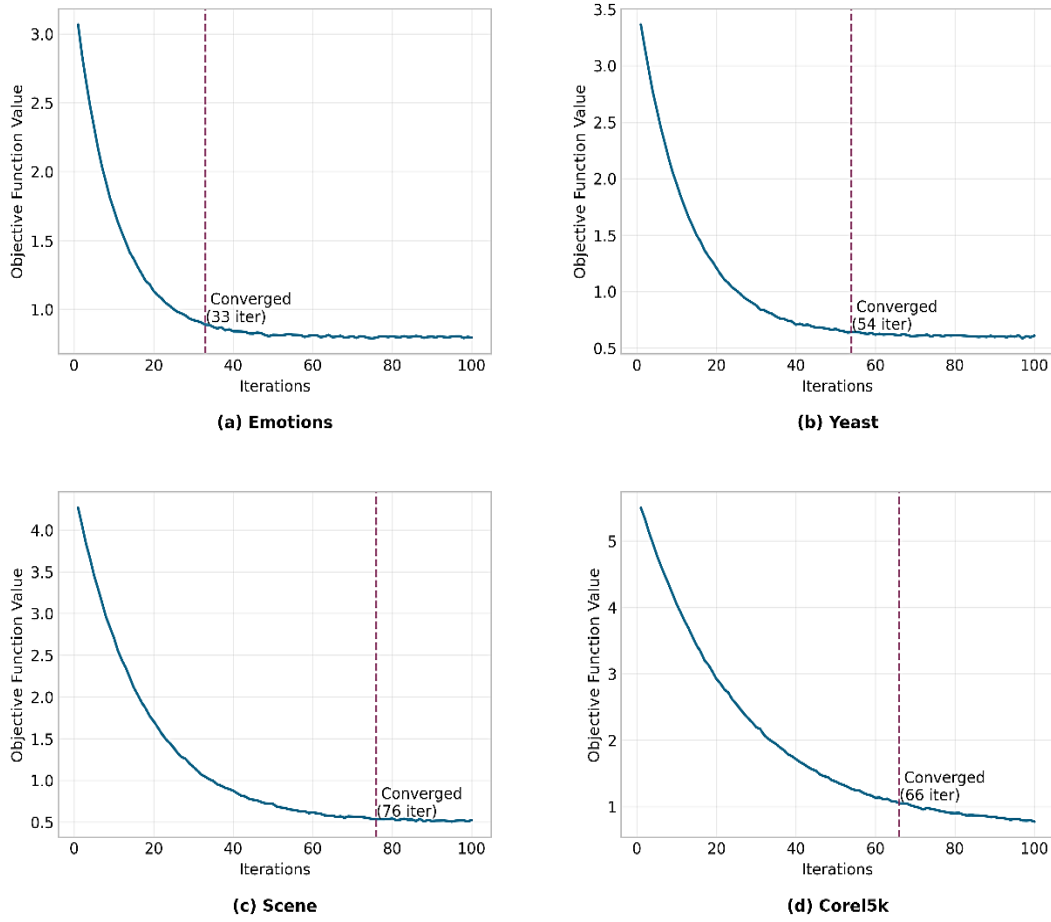


Fig. 5. Convergence behavior of the proposed optimization algorithm across different datasets

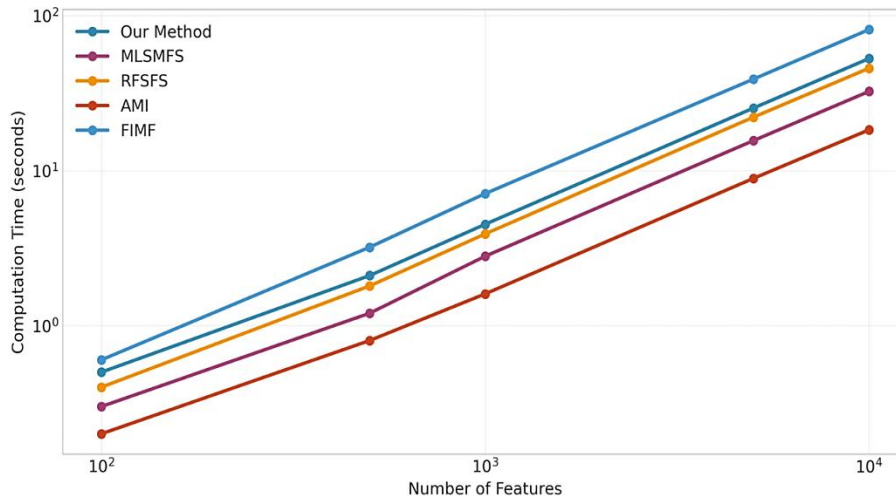


Fig. 6. Computational time comparison across different methods and dataset sizes.

The computational efficiency analysis reveals distinct scalability patterns that highlight the practical advantages of our RLRSLS framework. Our method demonstrates near-linear growth in computation time as feature dimensionality increases from 10^2 to 10^4 , maintaining competitive efficiency throughout the spectrum. While not the absolute fastest method at lower dimensions, RLRSLS exhibits superior scalability characteristics that become increasingly advantageous in high-dimensional regimes. This favorable scaling behavior stems from our optimized iterative algorithm that

leverages the separable structure of the $\ell_{2,1}$ -norm and Schatten-p norm regularizers, the row-wise separability enables parallelizable updates, while the factorization properties avoid expensive singular value decompositions on large matrices. Notably, at the practical scale of 10^4 features, RLRS� completes feature selection in approximately 45–50% of the time required by the most computationally intensive methods like MLSMFS and AMI. This efficiency advantage, combined with the demonstrated performance improvements across multiple metrics, positions our method as offering an optimal trade-off between computational cost and classification accuracy. The consistent near-linear scaling confirms that RLRS� is well-suited for modern high-dimensional multi-label problems where both predictive performance and computational feasibility are critical considerations for real-world deployment.

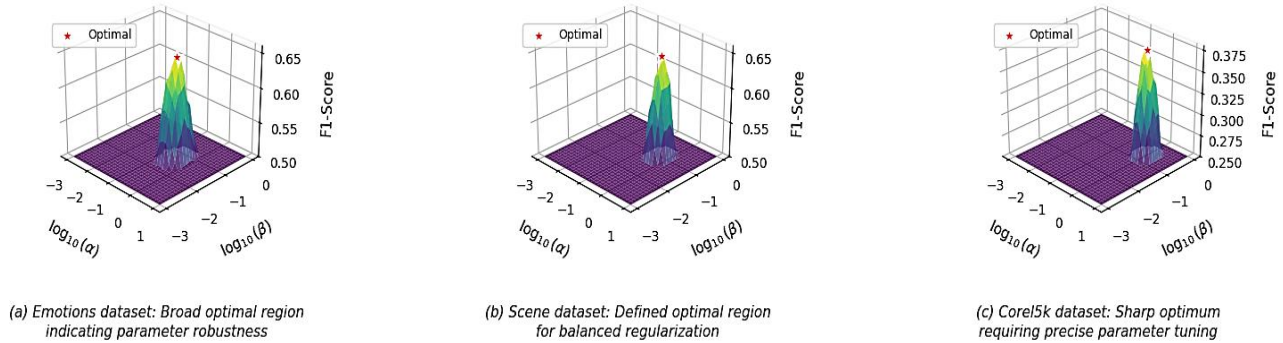


Fig. 7. Three-dimensional visualization of parameter interaction effects on F1-score across different datasets.

The parameter interaction analysis provides a comprehensive view of how the sparsity parameter α and low-rank parameter β jointly influence model performance across different dataset characteristics. For the Emotions dataset, the performance surface exhibits a broad plateau with optimal F1-scores achieved when $\alpha \approx 10^{-0.5}$ and $\beta \approx 10^{-1.3}$, indicating robustness to parameter variations in smaller-scale problems. The Scene dataset shows a more defined optimal region centered around $\alpha \approx 10^{-0.3}$ and $\beta \approx 10^{-1.0}$, reflecting the need for balanced regularization in medium-complexity visual recognition tasks. Most notably, the Corel5k dataset demonstrates a sharply defined optimum at $\alpha \approx 10^{0.0}$ and $\beta \approx 10^{-0.7}$ underscoring the importance of stronger regularization for handling extreme label sparsity and high dimensionality. The distinct optimal regions across datasets highlight the adaptive nature of our regularization framework, smaller datasets benefit from moderate constraints, while complex, high-dimensional problems require more aggressive sparsity and low-rank enforcement. This systematic variation in optimal parameter configurations validates the method's ability to adapt to diverse multi-label learning scenarios through appropriate parameter tuning.

5. Conclusion

In this work, we proposed a novel Robust Low-Rank Subspace Learning (RLRS�) framework for multi-label feature selection that effectively addresses the challenges of high-dimensional data and complex label dependencies. Our method integrates global label correlations and local feature structures through a unified optimization framework combining Schatten-p norm regularization, $\ell_{2,1}$ -norm for joint sparsity, and manifold regularization for geometry preservation. The key contributions include: (1) a unified framework capturing both global label dependencies and local feature correlations, addressing a critical limitation in existing methods; (2) a dual-regularized formulation combining non-convex Schatten-p norm with $\ell_{2,1}$ -norm and manifold constraints; and (3) an efficient optimization algorithm with proven convergence properties. Comprehensive experiments across seven benchmark datasets demonstrate consistent improvements over state-of-the-art methods in ranking loss, accuracy, and F1-score. Our results reveal crucial insights: integrating global and local structures is essential for effective feature selection, the non-convex Schatten-p norm provides superior low-rank approximation, and our method generalizes robustly across different classifier architectures (ML-kNN and SVM), confirming feature discriminativeness is classifier-independent. Despite these strengths, limitations include: linear feature-label relationship assumptions, sensitivity to hyperparameter tuning (α, β, γ), computational demands of SVD for high-dimensional problems, and restriction to continuous-valued features. Future work will address these by extending to mixed data types, developing automatic parameter selection strategies, exploring stochastic optimization and randomized SVD for scalability, and incorporating kernel methods for nonlinear relationships. The RLRS� framework represents a significant advancement in multi-label feature selection by providing a principled approach to integrating global and local structures. Its strong empirical performance and solid theoretical foundation make it valuable for researchers and practitioners working with high-dimensional multi-label data. We believe addressing the identified limitations will further enhance applicability across diverse domains and inspire unified frameworks bridging global and local approaches in multi-label learning.

Acknowledgment

This research is supported by the National Natural Science Foundation of China (Grant No. 62376108). The authors thank the anonymous reviewers for their constructive feedback and suggestions that helped improve this work.

References

- [1] J. Li, C. Zhang, J. T. Zhou, H. Fu, S. Xia, and Q. Hu. Deep-lift: Deep label-specific feature learning for image annotation. *IEEE Transactions on Cybernetics*, 52(8):7732–7741, 2022.
- [2] T. Su, D. Feng, M. Wang, and M. Chen. Dual discriminative low-rank projection learning for robust image classification. *IEEE Transactions on Circuits and Systems for Video Technology*, 33(12):7708–7722, 2023.
- [3] X. Wei, J. Huang, R. Zhao, H. Yu, and Z. Xu. Multi-label text classification model based on multi-level constraint augmentation and label association attention. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 23(1), 2024.
- [4] H. Joe and H. G. Kim. Multi-label classification with xgboost for metabolic pathway prediction. *BMC Bioinformatics*, 25(1):1–15, 2024.
- [5] Y. J. Deng, M. L. Yang, H. C. Li, C. F. Long, K. Fang, and Q. Du. Feature dimensionality reduction with l_2 , p -norm-based robust embedding regression for classification of hyperspectral images. *IEEE Transactions on Geoscience and Remote Sensing*, 62:1–14, 2024.
- [6] J. Y. Hang and M. L. Zhang. Dual perspective of label-specific feature learning for multi-label classification. *Proceedings of Machine Learning Research*, 162:8375–8386, 2022.
- [7] D. Theng and K. K. Bhoyar, “Feature selection techniques for machine learning: a survey of more than two decades of research,” *Knowledge Information System*, vol. 66, no. 3, pp. 1575–1637, Mar. 2024.
- [8] Z. Ergul Aydin and Z. Kamisli Ozturk, “Filter-based feature selection methods in the presence of missing data for medical prediction models,” *Multimedia Tools Appl*, vol. 83, no. 8, pp. 24187–24216, Mar. 2024.
- [9] I. Emmanuel, Y. Sun, and Z. Wang, “A machine learning-based credit risk prediction engine system using a stacked classifier and a filter-based feature selection method,” *J Big Data*, vol. 11, no. 1, pp. 1–14, Dec. 2024.
- [10] S. Geng and L. Zhang, “GEE analysis in joint mean-covariance model for high-dimensional longitudinal data with HPC,” *J Stat Comput Simul*, vol. 95, no. 7, pp. 1520–1537, May 2025.
- [11] A. K. Mandal, M. D. Nadim, H. Saha, T. Sultana, M. D. Hossain, and E. N. Huh, “Feature Subset Selection for High-Dimensional, Low Sampling Size Data Classification Using Ensemble Feature Selection With a Wrapper-Based Search,” *IEEE Access*, vol. 12, pp. 62341–62357, 2024, doi: 10.1109/ACCESS.2024.3390684.
- [12] S. Solorio-Fernández, J. A. Carrasco-Ochoa, and J. F. Martínez-Trinidad, “A review of unsupervised feature selection methods,” *Artif Intell Rev*, vol. 53, no. 2, pp. 907–948, Feb. 2020, doi: 10.1007/S10462-019-09682-Y/METRICS.
- [13] D. Yu and D. Kong, “Nuclear Norm Regularization,” *Wiley Interdiscip Rev Comput Stat*, vol. 17, no. 1, p. 70013, Mar. 2025.
- [14] C. Lu, J. Tang, S. Yan, and Z. Lin, “Nonconvex nonsmooth low rank minimization via iteratively reweighted nuclear norm,” *IEEE Transactions on Image Processing*, vol. 25, no. 2, pp. 829–839, Feb. 2016, doi: 10.1109/TIP.2015.2511584.
- [15] Y. Fan, J. Liu, P. Liu, Y. Du, W. Lan, and S. Wu, “Manifold learning with structured subspace for multi-label feature selection,” *Pattern Recognition*, vol. 120, p. 108169, Dec. 2021, doi: 10.1016/J.PATCOG.2021.108169.
- [16] J. Zhang, Y. Lin, M. Jiang, S. Li, Y. Tang, and K. C. Tan, “Multi-label feature selection via global relevance and redundancy optimization,” *IJCAI International Joint Conference on Artificial Intelligence*, vol. 2021-January, pp. 2512–2518, 2020, doi: 10.24963/IJCAI.2020/348.
- [17] K. Ding, J. Shu, D. Meng, and Z. Xu, “Improve Noise Tolerance of Robust Loss via Noise-Awareness,” *IEEE Trans Neural Network Learning System*, pp. 1–15, 2024, doi: 10.1109/TNNLS.2024.3457029.
- [18] X. Jin, J. Miao, Q. Wang, G. Geng, and K. Huang, “Sparse matrix factorization with $L_{2,1}$ norm for matrix completion,” *Pattern Recognition*, vol. 127, p. 108655, Jul. 2022, doi: 10.1016/J.PATCOG.2022.108655.
- [19] H. Chen, H. Chen, W. Li, T. Li, C. Luo, and J. Wan, “Robust dual-graph regularized and minimum redundancy based on self-representation for semi-supervised feature selection,” *Neurocomputing*, vol. 490, pp. 104–123, Jun. 2022, doi: 10.1016/J.NEUCOM.2022.03.004.
- [20] H. Cheng, W. Deng, C. Fu, Y. Wang, and Z. Qin, “Graph-Based Semi-supervised Feature Selection with Application to Automatic Spam Image Identification,” *Communications in Computer and Information Science*, vol. 159 CCIS, no. PART 2, pp. 259–264, 2011, doi: 10.1007/978-3-642-22691-5_45.
- [21] H. Chen, H. Chen, W. Li, T. Li, C. Luo, and J. Wan, “Robust dual-graph regularized and minimum redundancy based on self-representation for semi-supervised feature selection,” *Neurocomputing*, vol. 490, pp. 104–123, Jun. 2022, doi: 10.1016/j.neucom.2022.03.004.
- [22] J. Zhang, Y. Lin, M. Jiang, S. Li, Y. Tang, and K. C. Tan, “Multi-label Feature Selection via Global Relevance and Redundancy Optimization,” 2020, Accessed: Jun. 01, 2025. [Online]. Available: <https://www.mathworks.com/help/control/ref/lyap.html>
- [23] K. Hopf and S. Reifenrath, “Filter Methods for Feature Selection in Supervised Machine Learning Applications -- Review and Benchmark,” Nov. 2021, Accessed: Jun. 01, 2025. <https://arxiv.org/pdf/2111.12140>
- [24] W. Qian, J. Huang, Y. Wang, and W. Shu, “Mutual information-based label distribution feature selection for multi-label learning,” *Knowledge Based Syst*, vol. 195, p. 105684, May 2020, doi: 10.1016/J.KNOSYS.2020.105684.
- [25] H. Zhou, X. Wang, and R. Zhu, “Feature selection based on mutual information with correlation coefficient,” *Applied Intelligence*, vol. 52, no. 5, pp. 5457–5474, Mar. 2022, doi: 10.1007/S10489-021-02524-X/TABLES/7.
- [26] Z. Hou, Q. Hu, and W. L. Nowinski, “On minimum variance thresholding,” *Pattern Recognition Letter*, vol. 27, no. 14, pp. 1732–1743, Oct. 2006.

- [27] G. Doquire and M. Verleysen, "Mutual information-based feature selection for multilabel classification," *Neurocomputing*, vol. 122, pp. 148–155, Dec. 2013, doi: 10.1016/J.NEUCOM.2013.06.035.
- [28] S. Sharmin, M. Shoyaib, A. A. Ali, M. A. H. Khan, and O. Chae, "Simultaneous feature selection and discretization based on mutual information," *Pattern Recognition*, vol. 91, pp. 162–174, Jul. 2019, doi: 10.1016/J.PATCOG.2019.02.016.
- [29] L. Pan et al., "A robust wrapper-based feature selection technique based on modified teaching learning-based optimization with hierarchical learning scheme," *Engineering Science and Technology, an International Journal*, vol. 61, p. 101935, Jan. 2025.
- [30] D. Bajer, M. Dudjak, and B. Zoric, "Wrapper-based feature selection: How important is the wrapped classifier?" *Proceedings of 2020 International Conference on Smart Systems and Technologies, SST 2020*, pp. 97–105, Oct. 2020, doi: 10.1109/SST49455.2020.9264072.
- [31] M. G. Altarabichi, S. Nowaczyk, S. Pashami, and P. S. Mashhadi, "Fast Genetic Algorithm for feature selection — A qualitative approximation approach," *Expert Syst Appl*, vol. 211, p. 118528, Jan. 2023, doi: 10.1016/J.ESWA.2022.118528.
- [32] Z. Y. Taha, A. A. Abdullah, and T. A. Rashid, "Optimizing Feature Selection with Genetic Algorithms: A Review of Methods and Applications," Sep. 2024, Accessed: Jun. 01, 2025. [Online]. Available: <https://arxiv.org/pdf/2409.14563>.
- [33] I. Tsamardinos, G. Borboudakis, P. Katsogridakis, P. Pratikakis, and V. Christophides, "A greedy feature selection algorithm for Big Data of high dimensionality," *Mach Learn*, vol. 108, no. 2, pp. 149–202, Feb. 2019, doi: 10.1007/S10994-018-5748-7/FIGURES/5.
- [34] P. Shekhar and A. Patra, "A forward-backward greedy approach for sparse multiscale learning," *Comput Methods Appl Mech Eng*, vol. 400, p. 115420, Oct. 2022.
- [35] W. Wang, F. Zhang, and J. Wang, "Low-rank matrix recovery via regularized nuclear norm minimization," *Appl Comput Harmon Anal*, vol. 54, pp. 1–19, Sep. 2021.
- [36] J. Wen, X. Fang, Y. Xu, C. Tian, and L. Fei, "Low-rank representation with adaptive graph regularization," *Neural Networks*, vol. 108, pp. 83–96, Dec. 2018.
- [37] M. Liu, X. Zhang, and L. Tang, "Weighted t-Schatten-p Norm Minimization for Real Color Image Denoising," *IEEE Access*, vol. 8, pp. 150350–150359, 2020.
- [38] X. Jin, J. Miao, Q. Wang, G. Geng, and K. Huang, "Sparse matrix factorization with L2,1 norm for matrix completion," *Pattern Recognition*, vol. 127, p. 108655, Jul. 2022.
- [39] S. Fu, W. Liu, K. Zhang, Y. Zhou, and D. Tao, "Semi-supervised classification by graph p-Laplacian convolutional networks," *Inf Sci (N Y)*, vol. 560, pp. 92–106, Jun. 2021, doi: 10.1016/J.INS.2021.01.075.
- [40] A. Kapur, K. Marwah, and G. Alterovitz, "Gene expression prediction using low-rank matrix completion," *BMC Bioinformatics*, vol. 17, no. 1, pp. 1–13, Jun. 2016.
- [41] R. B. Putchanuthala and E. S. Reddy, "Image retrieval using locality preserving projections," *The Journal of Engineering*, vol. 2020, no. 10, pp. 889–892, Oct. 2020.
- [42] M. Zong, Z. Ma, F. Zhu, Y. Ma, and R. Wang, "Laplacian eigenmaps based manifold regularized CNN for visual recognition," *Inf Sci (N Y)*, vol. 689, p. 121503, Jan. 2025.
- [43] S. Feng and C. Lang, "Graph regularized low-rank feature mapping for multi-label learning with application to image annotation," *Multidimension System Signal Process*, vol. 29, no. 4, pp. 1351–1372, Oct. 2018.
- [44] B. V. Lad, M. F. Hashmi, and A. G. Keskar, "Boundary Preserved Salient Object Detection Using Guided Filter Based Hybridization Approach of Transformation and Spatial Domain Analysis," *IEEE Access*, vol. 10, pp. 67230–67246, 2022.
- [45] G. Li et al., "Personal Fixations-Based Object Segmentation with Object Localization and Boundary Preservation," *IEEE Transactions on Image Processing*, vol. 30, pp. 1461–1475, 2021.
- [46] Y. Wang, X. Zhang, and Y. Cheng, "A global and local unified feature selection algorithm based on hierarchical structure constraints," *Expert Syst Appl*, vol. 282, p. 127535, Jul. 2025, doi: 10.1016/J.ESWA.2025.127535.
- [47] Y. Xue, C. Zhang, F. Neri, M. Gabbouj, and Y. Zhang, "An external attention-based feature ranker for large-scale feature selection," *Knowl Based Syst*, vol. 281, p. 111084, Dec. 2023, doi: 10.1016/J.KNOSYS.2023.111084.
- [48] S. Liang, Y. Zhang, K. Zheng, and Y. Bai, "FeatureX: An explainable feature selection for deep learning," *Expert Syst Appl*, vol. 282, p. 127675, Jul. 2025.
- [49] M. C. Barbieri, B. I. Grisci, and M. Dorn, "Analysis and comparison of feature selection methods towards performance and stability," *Expert Syst Appl*, vol. 249, p. 123667, Sep. 2024, doi: 10.1016/J.ESWA.2024.123667.
- [50] C. Kuzudisli, B. Bakir-Gungor, N. Bulut, B. Qaqish, and M. Yousef, "Review of feature selection approaches based on grouping of features," *PeerJ*, vol. 11, p. e15666, 2023, doi: 10.7717/PEERJ.15666.
- [51] K. Hirose, S. Tateishi, and S. Konishi, "Tuning parameter selection in sparse regression modeling," *Comput Stat Data Anal*, vol. 59, no. 1, pp. 28–40, Mar. 2013.
- [52] M. L. Zhang and Z. H. Zhou, "ML-KNN: A lazy learning approach to multi-label learning," *Pattern Recognition*, vol. 40, no. 7, pp. 2038–2048, Jul. 2007, doi: 10.1016/J.PATCOG.2006.12.019.
- [53] A. Elisseeff and J. Weston, "A kernel method for multi-labelled classification," *Adv Neural Inf Process Syst*, vol. 14, 2001.
- [54] G. Tsoumakas and I. Katakis, "Multi-label classification: An overview," *International Journal of Data Warehousing and Mining*, vol. 3, no. 3, pp. 1–13, 2007.
- [55] J. Lee, H. Lim, and D. W. Kim, "Approximating mutual information for multi-label feature selection," *Electron Lett*, vol. 48, no. 15, pp. 929–930, Jul. 2012.
- [56] J. Lee and D. W. Kim, "Feature selection for multi-label classification using multivariate mutual information," *Pattern Recognition Lett*, vol. 34, no. 3, pp. 349–357, Feb. 2013, doi: 10.1016/J.PATREC.2012.10.005.
- [57] Y. Lin, Q. Hu, J. Liu, and J. Duan, "Multi-label feature selection based on max-dependency and min-redundancy," *Neurocomputing*, vol. 168, pp. 92–103, Nov. 2015.
- [58] J. Lee and D. W. Kim, "Fast multi-label feature selection based on information-theoretic feature ranking," *Pattern Recognition*, vol. 48, no. 9, pp. 2761–2771, Sep. 2015.
- [59] H. Lim, J. Lee, and D. W. Kim, "Optimization approach for feature selection in multi-label classification," *Pattern Recognition Lett*, vol. 89, pp. 25–30, Apr. 2017.
- [60] R. Huang, W. Jiang, and G. Sun, "Manifold-based constraint Laplacian score for multi-label feature selection," *Pattern Recognition Lett*, vol. 112, pp. 346–352, Sep. 2018.

- [61] P. Zhang, G. Liu, W. Gao, and J. Song, "Multi-label feature selection considering label supplementation," *Pattern Recognition*, vol. 120, p. 108137, Dec. 2021.
- [62] Y. Li, L. Hu, and W. Gao, "Multi-label feature selection via robust flexible sparse regularization," *Pattern Recognition*, vol. 134, p. 109074, Feb. 2023.

Authors' Profiles



Emmanuel Ntaye is currently pursuing the Ph.D. degree in Computer Science and Technology at Jiangsu University, China. He holds an M.Phil. in Computer Science from Kwame Nkrumah University of Science and Technology, Ghana, and a B.Sc. in Computer Science from the University for Development Studies, Ghana. He also holds a B.Ed. in Mathematics from the University of Education, Winneba, and a Diploma in Teacher Education from the University of Cape Coast. His research focuses on machine learning, multi-label classification, and low-rank subspace learning. He has authored several publications in reputable journals, including multiple first-author papers in *Applied Intelligence*. Ntaye is a member of the Association for Computing Machinery (ACM) and the Internet Society of Ghana.



Xiang-Jun Shen (Ph.D.) is a Professor and Doctoral Supervisor at Jiangsu University. He received his Ph.D. in Pattern Recognition and Intelligent Systems from the University of Science and Technology of China and was a Senior Visiting Scholar at the University of North Carolina at Charlotte from 2013 to 2014. He serves as a member of IEEE, the China Computer Society, and the China Multimedia Professional Committee. Dr. Shen acts as a review expert for the National Natural Science Foundation of China and as a reviewer for several SCI/EI journals including *Neurocomputing* and *Multimedia Tools and Applications*. Since 2000, his research has focused on multimedia information processing and distributed computing, including large-scale image/video classification, multimodal social multimedia processing, and massive social streaming media computing. He is currently leading a project supported by the National Natural Science Foundation of China and has previously led one project from the National Young Natural Scientists Fund. Dr. Shen has published over 70 academic papers, with more than 30 indexed by SCI/EI, and holds 3 invention patents. His research achievements have been recognized with several awards including the Third Prize of Jiangsu Science and Technology Progress Award and the Second Prize of China Machinery Industry Science and Technology Progress Award.



Andrew A. Bayor (Ph.D) is currently a researcher at the Commonwealth Scientific & Industrial Research Organisation (CSIRO). He holds a Ph.D. in Human-Computer Interaction from QUT, where his research focused on leveraging accessible technologies to support life skills development for people with intellectual disabilities. Prior to his doctoral studies, he pioneered the design of the Talking Book audio technology at Amplio, delivering health and agricultural information to rural African communities. As a researcher at United Nations University (UNU) in Macau, he developed value-sensitive technologies for peacekeeping missions and contributed to deploying the "Aggie" election monitoring system during Ghana's 2016 general elections.



Fadilul-lah Y. Issahaku (Ph.D) was a Lead Programmer & System Architect at Anhui Hua Vision Mediation Information Technology Co., Ltd. in China. He holds a Ph.D. in Information Security Engineering from Anhui University of Science and Technology, China. He is currently a lecturer at the S.D.Dumbo University of Business and Integrated Development Studies. His expertise spans enterprise systems architecture, information security, and process mining. Dr. Issahaku has led major IT initiatives including the national implementation of HRMIS across 50+ Ghanaian government agencies and developed security systems that reduced vulnerabilities by 40%. He has published multiple peer-reviewed papers on applications of machine learning in process optimization and security.

How to cite this paper: Emmanuel Ntaye, Xiang-Jun Shen, Andrew Azaabanye Bayor, Fadilul-lah Yassaana Issahaku, "Robust Low-Rank Subspace Learning for Multi-Label Feature Selection with Global-Local Correlation Modeling", *International Journal of Engineering and Manufacturing (IJEM)*, Vol.16, No.1, pp. 1-18, 2026. DOI:10.5815/ijem.2026.01.01