

Enhanced HSI classification using depth-wise separable CNN with advanced feature extraction and selection techniques

Yallamandaiah. S¹, M. Sanjana², N. Sindhu Bhargavi³, E. Brahmani⁴, S. Anjali⁵, T. Pravallika⁶

^{1,2,3,4,5,6} Department of Electronics and Communication Engineering, Vignan's Nirula Institute of Technology and Science for women, Guntur – 522002, Andhra Pradesh

yallamandaiah@gmail.com

ABSTRACT: Hyperspectral Image classification has many challenges. The challenges are caused due to the high dimensionality of the data and variability. The results in lower accuracy. The Existing methods cannot handle the high dimensional data, and poor feature selection and gives lower accuracy. To solve these issues this research addresses the ResNet50 model is used to extract the important features. IARO algorithm (Artificial Rabbits Optimization) is used to select the best features. The Classification is done using the Depth-wise Separable Convolutional Network. It is used to enhance the performance of the proposed method. The proposed model improves the quality of data and increases the accuracy of the model. This addresses the limitations of the traditional methods, and this proposed method offers more reliability and effective classification system. Overall, method performed better than many other existing approaches. It improved accuracy and also reduced the complexity of the process.

Keywords: Hyperspectral Image Classification, ResNet50, Feature Extraction, Artificial Rabbits Optimization (IARO), Depth-wise Separable CNN

1. Introduction

Hyperspectral Images are the advancements of the fastest advancements of the Satellite sensors. In Hyper Spectral image, each pixel has hundreds of spectral bands [1]. Hyper spectral images provide the detailed information [2-4]. They provide both spectral data and spatial data. The spectral data which gives the details about the materials. The spatial data which tells about the shape and the location. HSI classification goal is always been precise in the remote sensing research.

The HSI inaccurate categorization is caused due to two main variables. KNN is traditional classification method (K-Neighbour) used only spectral information of each pixel. SVM (Support Vector Machine) is also a traditional machine learning method used for classification that use only spectral information of each pixel. Spatial

information is ignored in the conventional classification methods and use only its spectral values for classification. Traditional methods look at only one pixel at a time and ignore the other pixels around it, so the accuracy of the results is reduced.

Convolution Neural Networks are used to classify the Hyper spectral image classification. CNN based methods improve classification accuracy using 2D and 3D convolution, batch normalization and ReLU. CNN has a limitation in detecting very small spectral differences between the neighbouring spectral bands. So, its performance becomes limited and less accurate.

2. Recent Works

This section discusses some previous studies related to Hyperspectral Image (HSI) analysis. To enhance superpixel representation using Multiple

Instance Learning (MIL), Huang et al. [7] proposed the Superpixel-based Multi-Scale Multi-Instance Learning (MSMIL) framework. This approach effectively captures both global relationships among superpixels and local information within individual superpixels. In addition, a spectral-spatial strategy was introduced to fuse multi-scale superpixel features.

To improve HSI classification performance, Zhang et al. [8] developed a Multi-Level Deformable Gated Aggregated (MDGA) network. The model integrates axis-decomposed convolution with deformable sampling to enable adaptive feature interaction across spatial locations. This mechanism improves the efficiency and sparsity of convolutional operations while reducing redundant information through gated feature selection. Furthermore, inverted residual blocks were incorporated into the network architecture to recover spatial features and reduce the model's overall complexity.

For wetland classification, Li et al. [9] introduced the Feature-Guided Dynamic Graph Convolutional Network (FGDGCN). This model includes a learnable superpixel generation module consisting of a superpixel generation block and a pixel-level feature enhancement block. The former improves feature discrimination, while the latter refines superpixel representations during training. In addition, a feature-guided adjacency matrix updating mechanism dynamically captures spectral-spatial relationships among nodes, enabling better aggregation of neighboring information. The learned features are then projected back to pixel space for wetland classification.

To address domain adaptation challenges, Xin et al. [10] proposed a Feature Disentanglement-Based Domain Adaptation Network (FDDAN). This framework separates domain-specific features from class-invariant features to generate domain-independent representations for classification tasks. A hybrid transformer-convolution architecture is used to extract both

local and global features. The model then applies two disentanglement processes to obtain domain-specific and domain-invariant attributes, enabling better alignment of domain distributions. Additionally, adversarial learning is employed to enhance the transferability and discriminative capability of the extracted domain-invariant features.

Finally, Liang et al. [11] introduced the Neighborhood Contrastive Tokenization (NeiCoT) method to generate compact and context-aware tokens for transformer-based encoding. The method maximizes shared information between local samples and their global anchor representation using a predictor applied to patch embeddings. This encourages meaningful feature learning and improves the relevance of nearby tokens. A token-level contrastive loss is further applied to align predictions with local samples while distinguishing them from other instances within a mini-batch. Gaussian weighting is also used to balance neighborhood contributions in the contrastive loss. The NeiCoT strategy is finally integrated into a sequence-based Masked Autoencoder (MAE) framework to obtain effective HSI representations.

3. Proposed Work Explanation

In hyperspectral image classification, proper preprocessing and effective feature extraction play a crucial role in improving classification results. This study proposes a comprehensive framework that incorporates advanced techniques such as noise removal, contrast improvement, and normalization. During the preprocessing stage, noise reduction methods are applied to eliminate sensor noise, contrast enhancement techniques are used to make the image clearer, and Min-Max normalization is implemented to standardize the pixel values across the hyperspectral data.

For feature extraction, an Improved ResNet50 model is utilized to effectively capture both spectral and spatial characteristics of the data with

higher precision. After extracting the features, an Improved Artificial Rabbits Optimization (IARO) algorithm is applied to select the most important features.

This optimization process helps in identifying the most relevant information, reducing unnecessary data, lowering computational cost, and enhancing overall model performance.

In the classification stage, an MRDSCNN model is used. This architecture combines multi-scale residual connections with depthwise separable convolution layers to accurately classify hyperspectral images at different spatial scales. By integrating these advanced methods, the proposed framework provides a reliable solution for hyperspectral image classification, achieving better accuracy while maintaining computational efficiency.

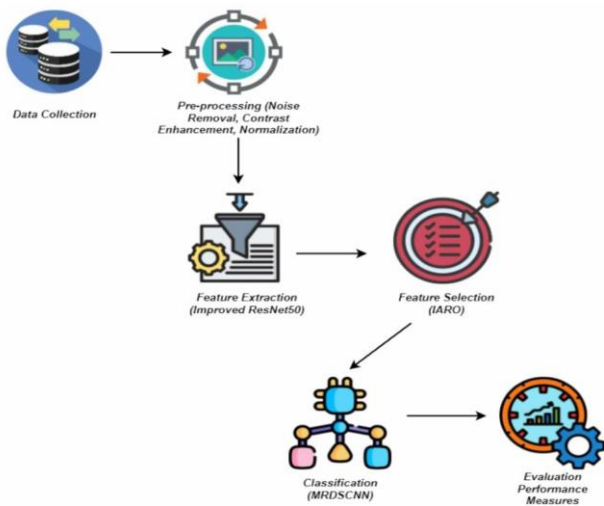


Figure 1: The Architecture of Proposed model

3.1 Data set description

This study uses three public datasets to test the method. The details of the three data sets are below:

Pavia University dataset [12]: It is collected by the ROSIS-03 Satellite over that area. It has 115 spectral bands, image size of 610 x 304, and 1.3m spatial resolution. After removing 12 noisy bands, the data is divided into nine classes for analysis.

Houston2013 Dataset [13]: This dataset was collected by ITRES CASI-1500 over the

University of Houston. It contains 349×1905 pixels with 144 spectral bands. The dataset includes 15 land cover classes such as roads, highways, and trees.

Xiongan Dataset [14]: This hyperspectral dataset has a size of 3750×1580 pixels with 250 spectral bands. It contains 20 types of land cover, including ash, pear, grassland, and residential areas.

3.2 Pre-processing

The pre-processing method includes the contrast enhancement, noise reduction, and min-max normalization all these provide the many benefits in the hyperspectral image categorization. To get the accurate output there are following stages in the pre-processing.

Noise Removal: To remove the noise in the image we use the mean filtering technique. The mean filtering technique and the smoothing of image stands for $g(x, y)$ and the (x, y) stands for the pixel point in an image with $f(x, y)$ where the neighbourhood 'S' having 'M' pixels in the following image.

$$G(x, y) = \frac{1}{M} \sum_{(i,j) \in S} f(x, y)(x, y) \notin S \quad (1)$$

Normalization: The Min-Max normalization represents the standardized original data, correctness, accelerate model convergence. It is calculated as follows:

$$x = \frac{x - \min(x)}{\max(x) - \min(x)} \quad (2)$$

To get detailed image as part of the normalization HIS's are encapsulated and aligned. Normalization has a source image. It is the technique for translating discrete subject-space data into a reference space.

Contrast Enhancement: CLAHE is the development of Histogram Equalization and Adaptive Histogram Equalization. CLAHE works more better than the Histogram Equalization and Adaptive Equalization. The CLAHE, improves the low-contrast images. Firstly, It improves the

brightness of each pixel in the image more quickly. It improves the contrast across the entire image. It is calculated by the following equation,

$$D_B = \frac{255}{8 \times 8} \sum_{i=0}^{D_A} H(i) \quad (3)$$

Secondly, it lessens the noise enhancement in the image. The final histogram equation:

$$H(i) = \begin{cases} H(i) + L, & H(i) < H_{max} \\ H_{max} + L, & H(i) \geq H_{max} \end{cases} \quad (4)$$

3.3 Feature extraction

After performing pre-processing the image, we further move to Feature extraction. In this we use the Improved ResNet50 technique to extract features from the image to perform the better classification. ResNet50 is a deep learning model and is used for feature extraction. It is used to skip connections which is to Avoid the loss of important information. It helps to train very deep networks easily. It can also capture the small and large features in images clearly. Hyperspectral images contain a lot of data includes many types of spectral Bands. Hyperspectral images can capture hundreds of narrow spectral bands for each pixel. Hyperspectral image helps to identifying materials, better classification accuracy, improves performance on benchmark datasets. Sub pixel information you can estimate what materials are mixed inside a single sub pixel. Hyperspectral images and deep networks work efficiently on large datasets. They are better for small object detection.

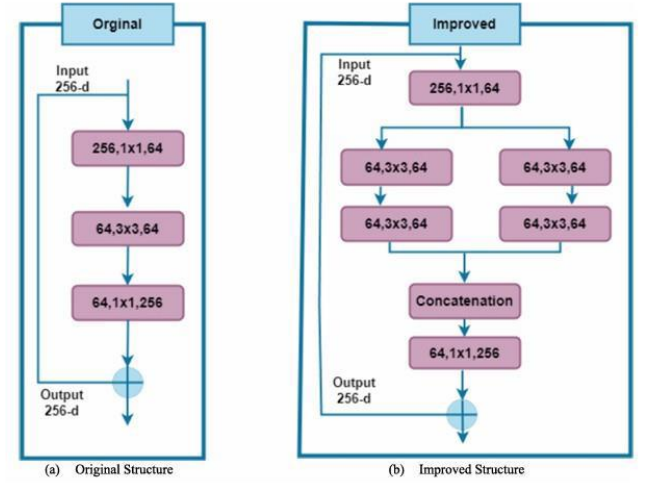


Figure 2: Comparison of the Original Structure and Improved Structure frame works

Table 1: The framework of RestNet50

Layer Name	50-Layer	Output Size
Conv1	$7 \times 7, 64, \text{stride}2$	112×112
Conv2_x	$3 \times 3 \text{ max pool, stride}2$	56×56
Conv3_x	$\begin{bmatrix} 1 \times 1, 64 \\ 3 \times 3, 64 \\ 1 \times 1, 256 \end{bmatrix} \times 3$	28×28
	$\begin{bmatrix} 1 \times 1, 128 \\ 3 \times 3, 128 \\ 1 \times 1, 512 \end{bmatrix} \times 4$	
	$\begin{bmatrix} 1 \times 1, 256 \\ 3 \times 3, 256 \\ 1 \times 1, 1024 \end{bmatrix} \times 6$	
Conv4_x	$\begin{bmatrix} 1 \times 1, 512 \\ 3 \times 3, 512 \\ 1 \times 1, 2048 \end{bmatrix} \times 3$	7×7
Conv5_x	Average pool, 1000-d fc, softmax	1×1

HSI data also improve the performance of the detection of models like faster-CNN. Original structure is a single path which means the data flows through me sequence of layers. It is having the basic

feature extraction. It is in same filter size everywhere. It can't capture small and large patterns at the same depth. Original structure works well only for simple datasets.

Improved model having multiple paths each path can use different type of kernel sizes or operations. Improved model having advanced feature extraction. It is having more detailed and accurate. It captures more details. It looks at data from different angles. It helps to learn complex patterns.

And gives better classification results works better where data is very complete.

3.4 Feature Selection

IARO means Improved Artificial Rabbit Optimization Algorithm. This IARO is inspired by the behaviour of rabbits in nature. Rabbits search for food smartly and choose safe places to hide from danger. Similarly, this algorithm helps to select the best data features. IARO is used to select the important features only. It is used to remove unnecessary data. It is used to improve accuracy and speed. There are few stages in this algorithm. Firstly, Random Hiding Strategy. In these rabbits digging several burrows for hiding and to stay safe. Similarly, the algorithm tries different feature combinations or better representation. It increases the variety. And avoid getting stuck in one solution.

Secondly, detour foraging strategy here rabbits go far away to search for food and explores more possibilities. They don't get stuck using only the same local features. Making multiple paths or global search helps to avoid wrong best solution. After searching everywhere, it is easy to pick the best features.

Thirdly, Energy conservation strategy rabbits save their energy reducing their movement. The algorithm starts with exploration then slowly focuses on the best solution. It keeps only useful data. It removes noise and unnecessary data. It helps to Improve accuracy. It helps for lower computational cost. Slowly moving from exploration to get best solution to make results stable and accurate.

3.5 Classification

DSCNN is the Depth wise separable convolutional Neural Network. It is used for the classification of hyperspectral images. HSI contains more than hundred bands per pixels. It is very huge and messy. DSCNN takes all the data without clashing or getting confused. It gives better accuracy and less computation cost. Multi scale feature extraction models looks at both tiny details like

edges and large patterns like full regions together at the same time.

Residual learning uses skip connections. They add the input directly to the output. It solves the vanishing gradient problem to makes training easier. Depth wise separable convolution, normal convolutions are heavy computation compared to depth wise separable convolution. It first applies to each band separately to capture the spatial features.

First, we have to learn each feature and then combine them smartly. The CNN architecture, the size of input image $80 \times 6 \times 3$. Two depth wise separable convolution blocks are applied. Batch Normalization for normalize data and add non-linearity (ReLU). Channels attention module focuses on important features and ignores unwanted ones. Squeeze step compressing the information using average pooling. Excitation learns which channels are important. Scale gives more importance to useful features. Skip connection improves performance by adding original input to final output. Output feature map size is $80 \times 6 \times 3$. It is same size as input. And better feature representation with low computational cost.

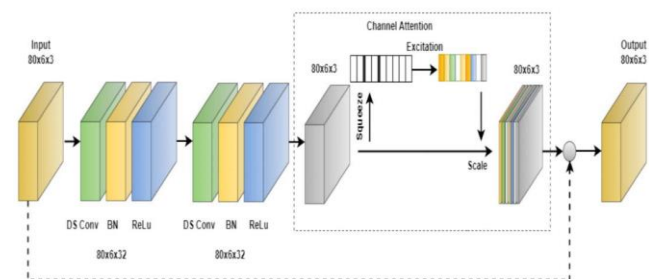


Figure 3: Architecture of CNN

4. Experimental results

Every experiment is done using Windows having intel core i9-10900 K CPU and 128 GB RAM, and NVIDIA GeForce RTX 3090 GPU (24 GB GDDR6X). The performance of the data sets is here:

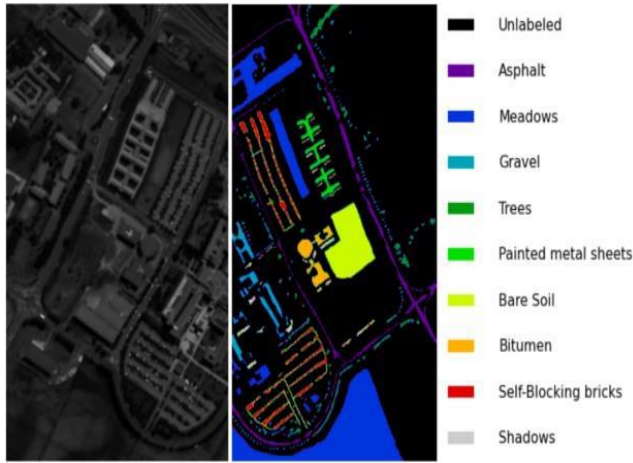


Figure 4: Pavia University dataset

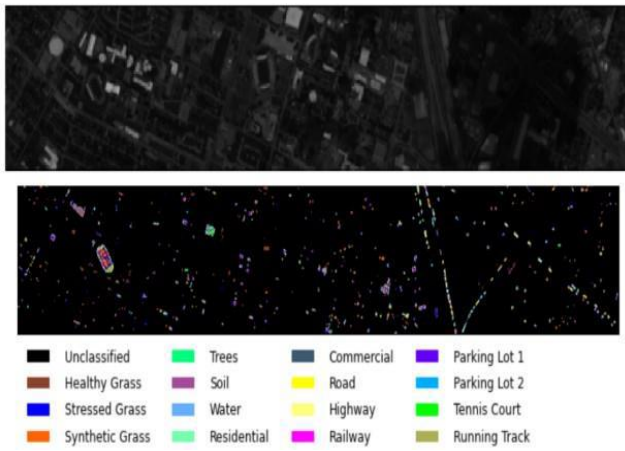


Figure 5: Houston 2013 dataset



Figure 6: Xiongan dataset

4.1 The performance of the Pavia University Dataset

The table 2 represents the categorized results of Pavia university dataset. And the fig 4a represents the Pavia university dataset description. It represents the false color image, label information and ground truth map of the image form the top to bottom. The proposed method constant results best classification accuracy over 9 classes shown in the table 2 in which many categories resulting a score of 100%. In that table the class 5 and 9 are having less difference in accuracy, compared to high scores. It performs well particularly and we can observe in class 6 and 7. It's (Overall Accuracy) OA of 99.89%, (Average Accuracy) AA of 99.98%, and (Kappa Accuracy) KA of 99.96%. It offers informative gain in classification accuracy.

Table 2: Categorization results of the Pavia University Dataset.

Class	SVM	KNN	2D-CNN	3D-CNN	ViT	Deep ViT	CaiT	CvT	Proposed
1	88.54	92.33	96.69	96.59	96.46	96.49	94.92	96.52	100
2	94.24	94.33	92.74	92.61	92.64	92.75	91.79	92.70	100
3	65.72	77.04	94.33	94.02	93.82	94.76	88.70	94.35	100
4	93.98	92.80	97.88	97.58	97.54	97.79	97.14	97.92	100
5	99.52	99.47	100	99.98	99.98	99.99	99.90	99.99	99.96
6	77.19	79.62	100	99.96	99.85	99.97	96.00	100	100
7	53.96	82.80	100	99.65	99.62	99.93	95.48	100	99.89
8	84.32	84.73	100	99.78	99.50	99.91	97.30	99.99	100
9	100	99.98	100	99.86	99.95	99.96	99.89	99.96	100
AA	84.16	89.23	97.96	97.78	97.71	97.95	95.68	97.94	99.9833
OA	88.96	90.58	92.29	92.09	92.07	92.27	90.60	92.21	99.98
KA	85.02	87.36	90.14	89.90	89.87	90.12	87.97	90.05	99.96

4.2 The performance of the Houston 2013 Dataset

The table 3 represents categorization outcomes of the Houston 2013 dataset. This uses conventional techniques like KNN and SVM, having lesser accuracy. The lesser accuracy is particularly shown in class 9 and 13. SVM struggle with 70.07% and 12.69% accuracy. KNN does not match with deep learning model's output. 3D-CNN performs well in class 3 and 4 with 100% accuracy. 3D-CNN overall accuracy is 98.85%. it achieves perfect accuracy in most classes. ViT variations are Deep ViT, CaiT and CvT with OA of 99.59%. These manage complex spectral-spatial interactions skillfully. This method's AA of 99.85%, OA of 99.83% and KA of 99.83%. It shows lesser accuracy than PaviaU dataset.

Table 3: Categorization results of the Pavia University Dataset.

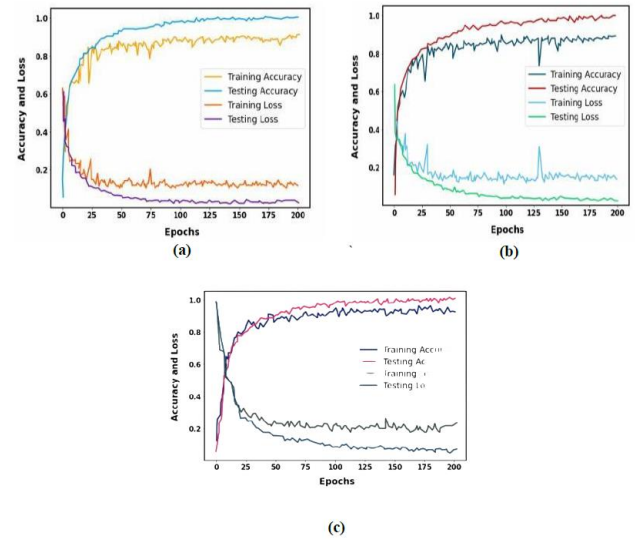
Class	SVM	KNN	2D-CNN	3D-CNN	ViT	Deep ViT	CaiT	CvT	Proposed
1	92.43	98.88	96.20	99.60	99.45	99.73	98.73	99.71	100
2	94.19	98.90	98.74	99.33	99.17	99.43	98.92	99.51	100
3	99.06	99.62	99.84	100	99.95	100	99.83	100	100
4	95.88	99.53	96.89	100	99.67	99.77	99.04	99.93	99.60
5	94.76	98.19	98.49	98.51	98.49	98.69	97.89	98.53	100
6	90.71	99.33	94.09	98.98	98.21	99.85	98.86	100	100
7	74.45	95.04	94.93	98.94	98.62	99.34	98.75	99.42	99.86
8	70.56	95.15	92.11	99.19	98.61	99.12	97.81	99.14	99.47
9	70.07	88.06	92.96	98.53	98.93	99.22	97.51	99.09	100
10	65.81	91.05	91.47	99.59	99.47	99.86	97.46	99.96	99.46
11	66.60	90.37	95.05	98.09	98.60	98.85	97.38	98.68	99.73
12	61.26	89.06	93.87	99.73	99.20	99.85	97.39	99.97	100
13	12.69	64.09	95.68	99.29	98.39	99.32	94.33	99.89	100
14	86.85	98.69	99.28	100	99.69	99.96	99.54	100	100
15	99.04	99.47	98.79	100	99.87	100	99.40	100	100
AA	78.29	93.70	95.89	99.32	99.09	99.53	98.19	99.59	99.85
OA	79.78	94.35	95.16	98.85	98.85	99.07	97.72	99.02	99.83
KA	78.11	93.89	94.77	98.76	98.76	99.00	97.53	98.94	99.83

4.3 The performance of the Xiongan Dataset

The table 4 represents Xiongan dataset categorization. In the fig 4c we can find the false color map and ground truth of the Xiongan dataset. When compared to 2D-CNN and ViT the performance model performance is so better. The proposed method is having the high accuracy rate compared to those traditional methods and all. It's Overall Accuracy is 99.59%, AA is 99.75%, and (Kappa Coefficient) KA is 99.64%.

Table 4: Categorization results of the Xiongan Dataset

Class	KNN	2D-CNN	3D-CNN	ViT	Deep ViT	CaiT	CvT	Proposed
1	74.69	97.09	92.21	98.73	98.08	47.31	99.08	99.98
2	73.90	99.11	96.75	99.80	99.40	57.00	99.96	99.89
3	76.98	97.23	97.03	99.71	99.09	52.93	99.91	99.12
4	97.91	99.36	99.25	99.49	99.46	94.41	99.51	99.51
5	69.87	98.77	95.50	99.51	98.87	54.54	99.92	99.71
6	72.26	99.06	95.51	99.26	98.80	48.96	99.70	100
7	70.90	99.73	99.70	99.95	99.83	26.62	99.99	99.96
8	94.97	99.47	99.35	99.82	99.70	89.09	99.85	100
9	96.33	99.60	99.62	99.95	99.82	93.14	99.97	99.86
10	95.69	99.88	99.70	99.98	99.94	89.34	99.98	100
11	6.64	91.95	76.07	97.77	95.57	0.00	99.46	100
12	59.96	95.20	89.54	98.97	97.91	27.13	99.53	99.75
13	77.64	98.31	95.21	99.32	98.83	74.21	99.73	99.91
14	56.91	91.69	85.27	98.60	97.78	4.54	99.56	99.77
15	62.78	94.87	87.74	98.88	98.35	27.78	98.94	99.47
16	36.77	84.23	68.96	97.32	95.80	21.69	99.35	99.41
17	2.51	53.26	39.98	91.94	90.10	0.07	98.92	99.64
18	74.07	97.81	93.87	99.25	98.66	62.46	99.49	99.89
19	56.79	97.80	91.73	99.24	98.58	54.14	99.90	99.23
20	90.55	96.75	96.92	99.01	98.58	83.57	99.19	99.34
AA	67.41	94.56	90.00	98.83	98.16	50.45	99.60	99.75
OA	78.82	98.08	95.23	99.15	98.70	68.31	99.44	99.59
KA	75.32	97.77	94.46	99.01	98.50	62.90	99.36	99.64

**Figure 7: Performance of evaluation of proposed model accuracy a. Pavia University dataset, b. Houston 2013 dataset, c. Xiongan dataset**

4.4 The overall performance of the proposed method

It's overall performance is evaluated on the basis of the different types of accuracy we considered. Based on that we evaluate the Pavia University dataset shows more accuracy in many different classes considered in the table 2, 3 and 4. In the proposed method the accuracy is calculated using the OA, KA and AA. During the training the model's performance is evaluated to minimize the errors. In testing, the model HSI classification is evaluated and to check the robustness and generalization capabilities.

5. Conclusion

In conclusion, HSI is challenging because it involves very high dimensional data and a lot of variation. These factors make accurate and efficient processing difficult. Traditional approaches often struggle because they cannot effectively handle such complex, limited generalization capabilities, high dimensional data and inefficiencies in feature selection. The study proposes a new method that overcomes limitations of earlier approaches by using contrast enhancement, noise reduction and normalization, leading to better accuracy and efficiency. The method improves data quality through

preprocessing and uses an enhanced ResNet50 for feature extraction along with the IARO algorithm to select key features and reduce complexity. A multi-scale CNN with residual and depth wise separable convolutions boosts classification performance. Overall, it achieves high accuracy with lower computational cost and performs well even with limited data.

References

1. Li Z, Zheng K, Li J, Li C, Gao L (2024) Cross semantic heterogeneous modeling network for hyperspectral image classification. *IEEE Trans GeosciRemoteSens*. <https://doi.org/10.1109/TGRS.2024.3426358>
2. Xu J, Li K, Li Z, Chong Q, Xing H, Xing Q, Ni M (2024) Fuzzy graph convolutional network for hyperspectral image classification. *Eng Appl Artif Intell* 127:107280. <https://doi.org/10.1016/j.engappai.2023.107280>
3. Jiang M, Su Y, Gao L, Plaza A, Zhao XL, Sun X, Liu G (2024) GraphGST: Graph generative structure-aware transformer for hyperspectral image classification. *IEEE Trans Geosci RemoteSens*. <https://doi.org/10.1109/TGRS.2023.3349076>
4. Ahmad M, Usama M, Khan AM, Distefano S, Altuwaijri HA, Mazzara M (2024) Spatial spectral transformer with conditional position encoding for hyperspectral image classification. *IEEE Geosci RemoteSensLett*. <https://doi.org/10.1109/LGRS.2024.3431188>
5. Huang Y, Peng J, Zhang G, Sun W, Chen N, Du Q (2024) Adversarial domain adaptation network with calibrated prototype and dynamic instance convolution for hyperspectral image classification. *IEEE Trans Geosci Remote Sens*. <https://doi.org/10.1109/TGRS.2024.3387990>
6. Lv H, Li Y, Zhang H, Wang R (2024) Multi-dimensional deep dense residual networks and multiple kernel learning for hyper spectral image classification. *Infrared Phys Technol* 138:105265. <https://doi.org/10.1016/j.infrared.2024.105265>
7. Huang S, Liu Z, Jin W, Mu Y (2024) Superpixel-based multi scale multi-instance learning for hyperspectral image classification. *PatternRecogn*149:110257. <https://doi.org/10.1016/j.patcog.2024.110257>
8. Zhang Z, Zhou H, Zhang C, Zhang X, Jiang Y (2023) A multi-level deformable gated aggregated network for hyperspectral image classification. *Int J Appl Earth Obs Geoinf* 123:103482. <https://doi.org/10.1016/j.jag.2023.103482>
9. Li Z, Meng Q, Guo F, Wang L, Huang W, Hu Y, Liang J (2023) Feature-guided dynamic graph convolutional network for wet land hyperspectral image classification. *Int J Appl Earth Obs Geoinf*123:103485. <https://doi.org/10.1016/j.jag.2023.103485>
10. Xin Z, Li Z, Xu M, Wang L, Ren G, Wang J, Hu Y (2024) Feature disentanglement-based domain adaptation network for cross-scene coastal wetland hyperspectral image classification. *Int J Appl Earth Obs Geoinf*129:103850. <https://doi.org/10.1016/j.jag.2024.103850>
11. Liang M, Zhang X, Yu X, Yu L, Meng Z, Zhang X, Jiao L (2024) An efficient Transformer with neighborhood contrastive tokenization for hyperspectral images classification. *Int J Appl EarthObsGeoinf*131:103979. <https://doi.org/10.1016/j.jag.2024.103979>
12. Kavitha K, Arivazhagan S, Suriya B (2014) Classification of Pavia University hyperspectral image using Gabor and SVM classifier. *Int J New Trends Electron Commun* 2(3): 9–14.
13. Liu H, Li W, Xia X G, Zhang M, Tao R (2023) Multiarea target attention for hyperspectral image

classification. IEEE Trans Geosci Remote Sens 61:

<https://doi.org/10.1109/TGRS.2023.3318071>

14. Luo X, Tong X, Pan H (2020) Integrating multi-resolution and multitemporal Sentinel-2 imagery for land-cover mapping in the Xiongan New Area, China. IEEE Trans Geosci Remote Sens 59(2):10291040. <https://doi.org/10.1109/TGRS.2020.2999558>