

Classification of peanut variety based on hyperspectral imaging and improved extreme learning machine

MENGKE WANG, HONGRUI ZHANG, HONGFEI LV, CHENGYE LIU, JINHUAN XU*,
XIANGDONG LI

Institute of Automation, Qilu University of Technology (Shandong Academy of Sciences), Jinan, China

**Corresponding author: jhxu@sdas.org*

Citation: Wang M., Zhang H., Lv H., Liu C., Xu J., Li X. (2025): Classification of peanut variety based on hyperspectral imaging and improved extreme learning machine. Czech J. Food Sci., 43: 17–28.

Abstract: Peanut as an important crop, plays an important role in agricultural production, which is rich in edible vegetable oil and protein. The variety of peanut affects the content of vegetable oil and protein. Therefore, the classification of peanut variety can better promote the sustainable development of agriculture. In this study, hyperspectral imaging technology is used to achieve peanut variety classification. In addition, the spatial-spectral extreme learning machine (SS-ELM) is proposed to process the hyperspectral data to get the final classification label. To fully explore the spatial structure information of hyperspectral data, propagation filtering is integrated into the framework of extreme learning machine (ELM). The average accuracy of the improved ELM model on five varieties of peanuts dataset (Luhua 11, Dabaisha, Xiaobaisha, Fenghua, and Luohanguo 308) is 98.32%, which is higher than other classic models. The experimental results show that the improved ELM can classify peanut of different varieties by hyperspectral imaging.

Keywords: hyperspectral technology; machine learning; propagation filtering; spatial-spectral information

Peanut, also known as peanut and evergreen fruit, is an important edible oil raw material and cash crop in more than 100 countries around the world. Currently, China is the world's largest producer of peanuts, and its output value ranks fourth among major crops, after rice, corn and wheat. Peanut has the advantages of large-scale production, efficient planting, high efficiency of oil production, excellent oil quality, strong international competitiveness and broad demand prospects (Carrin and Carelli 2010). It has obvious advantages and potential in ensuring the supply of edible

vegetable oil in China. In recent years, in order to meet the increasing demand of agriculture and industry (Wright 1980), seed hybridisation technology has been widely used. However, this has also led to a rapid increase in peanut varieties, with more than 500 peanut varieties now available, including more than 30 excellent varieties. The cultivation of excellent peanut varieties is very important for peanut cultivation.

Although there are many varieties of peanuts, the differences between them are sometimes difficult to distinguish by the naked eye, and even the peanut

Supported in part by the Pilot Special Basic Research Project for Science, Education, and Industry Integration, under Grant No. 2023PX061; in part by the Talent Scientific Research Project of Universities (Institutes) under Grant No. 2023RCKY156; in part by Shandong Province key research and development plan under Grant No. 2024CXGC010210; in part by the Key research and development program of Shandong Province Nr. 2021JMRH0108; in part by Shandong Provincial Natural Science Foundation under Grant No. ZR2022MF312.

of the same variety may vary in particle size. In addition, there are also various situations that lead to the mixing of different varieties during peanut growth, such as during harvesting, transportation and storage (Liu et al. 2022). These factors result in the complexity and diversity of peanut varieties, and also increase the difficulty of peanut quality control. Different varieties of peanuts have different nutritional composition and characteristics, and the different planting environment and way will also affect the growth of peanuts. Mixing different varieties of peanuts may lead to lower yields, so it is important to identify peanut varieties before planting.

In the process of classifying peanut varieties, the traditional method of variety classification is manual identification, where the identification relies on the human eye to observe the shape of the peanut seed, peel characteristics, colour, and other appearance features, which requires rich experience to accurately identify different varieties of peanuts. However, this method has disadvantages such as destructiveness, time-consuming and laborious, limited accuracy, and low efficiency (Fabiya et al. 2020). To solve the above drawbacks of traditional identification methods, non-destructive testing technology has been proposed, machine vision (Mohi-Alden et al. 2023), near infrared spectroscopy (Jiang et al. 2022), hyperspectral imaging technology have been widely used in agriculture and food inspection.

Machine vision is mainly based on morphological characteristics, and has been widely used in the breakage detection of seeds. But for certain different varieties peanuts with similar colour and same morphology (Wang et al. 2015), they can no longer be accurately and efficiently classified by machine vision.

Near-infrared spectroscopy is a well-established non-destructive testing technique that can be used for non-destructive testing of complex samples (Zhang et al. 2022). NIR spectroscopy is used to obtain spectral information about a substance by measuring the absorption or reflection of near-infrared (NIR) light, and is now widely used to analyse the internal chemical composition of food (Qu et al. 2015). Although NIR has a wide spectral range and high spectral resolution, NIR spectral data are presented in the form of spectral curves, but do not provide spatial information about the sample, such as shape, size, and location. If the spatial information of the sample composition is not considered when analysing the data, a large amount of important information may be lost, affecting the processing results.

Hyperspectral imaging is an emerging technology that incorporates traditional imaging techniques and spectroscopy. It acquires both spatial and spectral information about an object (Huan et al. 2019), and the data are presented as a three-dimensional spatial cube ($x \times y \times \lambda$), where x and y denote the spatial dimensions, λ denotes the continuous wavelength. Each pixel corresponds to a spectral curve that reflects the characteristic reflectance properties. Thus, spectral information can uniquely describe, and distinguish substance types (Elmasry et al. 2012).

In recent years, with the continuous development of hyperspectral imaging technology, it has now become a research hotspot for agricultural product detection. Yuan et al. (2020) studied the mouldy peanuts recognition by a small number by critical band cohesion classifiers (EC) based on hyperspectral images. Then, support vector machine (SVM)-based EC, partial least squares-discriminant analysis (PLS-DA), and cluster-independent soft pattern classifiers (SIMCA) were used to select the critical wavelengths to identify healthy and mouldy peanuts. The overall pixel classification accuracies of EC, SVM, PLS-DA and SIMCA are 97.66, 97.53, 95.31, and 97.36%, respectively. Liu et al. (2020) established Deeplab v3+, Segnet, Unet, and Hypernet as control models, integrated peanut Identification Index (PRI) into hyperspectral images as pre-extraction of data features, and then integrated the constructed multi-feature fusion block into the control model as a feature enhancement module. The experimental results show that the average accuracy of four control models is increased by 0.61% to 1.15%. Through feature enhancement, the accuracy of the model is improved by 0.43–4.96%. This indicates that the proposed method has a significant effect in improving the accuracy of peanut recognition. Zou et al. (2022) studied the classification algorithms of different peanut varieties based on hyperspectral imaging technology, and the best classification algorithm among them was MF-LightGBM-LightGBM-Optuna-LightGBM, after using various data preprocessing methods and feature band extraction methods. The optimisation effect of the model was obvious after using Optuna algorithm to optimise the model, the optimisation effect of LightGBM is obvious, especially in running time, which is 11 times faster than XGBoost and 16 times faster than before optimisation. Wu et al. (2022) pre-processed the hyperspectral data with median filtering and adopted four variable selection methods to obtain characteristic wavelengths for mildewing detection of peanuts. The study compared the performance

<https://doi.org/10.17221/109/2024-CJFS>

of stacked ensemble learning (SEL) model with extreme gradient boosting algorithm (XGBoost), illumination gradient boosting algorithm (LightGBM), and type boosting algorithm (CatBoost). The results show that MF-LightGBM-SEL model has the highest prediction accuracy, reaching 98.03%, and the modelling time is only 0.37 s.

Existing classification methods in processing hyperspectral data have problems such as high computational cost and sensitivity to data noise, and hyperspectral data itself is characterised by rich information and high computational complexity. Therefore, to overcome these challenges and make more effective use of hyperspectral data, a classification method based on a space-spectral extreme learning machine for the rapid and lossless classification of peanut varieties is proposed to improve classification accuracy, reduce computational costs, and enhance model stability and generalisation ability. The main contribution of this article are as follows: *i*) The hyperspectral imaging technology is used to obtain the spectral information under different bands, which effectively solves the problem of same-colour foreign matter. *ii*) The spatial-spectral extreme learning machine (SS-ELM) model is proposed by combining spatial and spectral information into the extreme learning machine (ELM) framework, and to achieve peanut variety classification by using spatial filtering that incorporates local spectral spatial context integration and reshaping mechanisms into the hidden layer feature representation.

The rest of this article is organised as follows. The second part briefly introduces sample collection, label preparation and ELM algorithm. In section Material and Methods, the SS-ELM algorithm for hyperspectral image classification is proposed, and the method of hyperspectral image classification by spatial information propagation filtering fusion is introduced in detail. The experimental results are given in section Results and Discussion. Finally, section Conclusion summarises the thesis.

MATERIAL AND METHODS

Sample preparation

Five peanut varieties (Luhua 11, Dabaisha, Xiaobai-sha, Fenghua, and Luohangguo 308) with similar appearance are taken as research objects. Five peanuts from each variety were selected as experimental samples. The sample was placed flat on a 17.5 × 17.5 cm black tray and the data set was collected using a spectral camera.

Hyperspectral image acquisition

In the experiment of peanut hyperspectral image acquisition, a near-infrared hyperspectral camera of the 'FS-15' series from Hangzhou Caipu company and fig-spec software were used. The effective spectral band range of this camera is 900–1 700 nm, and the spectral resolution is 6 nm. There are 256 bands in total, and the pixel size of the imaging is 601 × 320.

To capture the images, the peanut sample tray was placed on a motorised stage as shown in Figure 1A–B. The distance between the peanut sample and the camera lens was set to 0.26 m, and the moving speed of the sample was set to 20.801 mm·s⁻¹. The exposure time of the hyperspectral camera was set as 15 000 μs. Due to the influence of noise caused by environment factors and dark current of the instrument, it is necessary to execute black and white correction respectively before sample collection.

Spectral correction

Because the dark current generated by the thermal movement inside the electronic components will affect the signal-to-noise ratio of the image, it is necessary to compensate the influence of the dark current on the data through dark current correction to improve the quality and stability of the data. And the spectral distribution of the light source may not be uniform, the light intensity at each wavelength is also different, so it is also necessary to eliminate the difference in spectral response at different wavelengths by white correction. The correction formula is:

$$R_c = \frac{R_0 - R_B}{R_w - R_B} \quad (1)$$

where: R_c – image after black and white correction; R_0 – original hyperspectral image; R_B – black correction image with the lens cap on; R_w – white correction plate image.

Label preparation

In order to evaluate the effectiveness of the algorithm, a corresponding set of labels is made, as shown in Figure 2. The threshold segmentation method is used to binarise the image, and a 3 × 3 mean filter (Gupta 2011) is used to eliminate most of the noise in the background, then the binary map is converted to a gray-scale map. And then a Gaussian filter with a standard deviation of 1 (Deng and Cahill 1993)

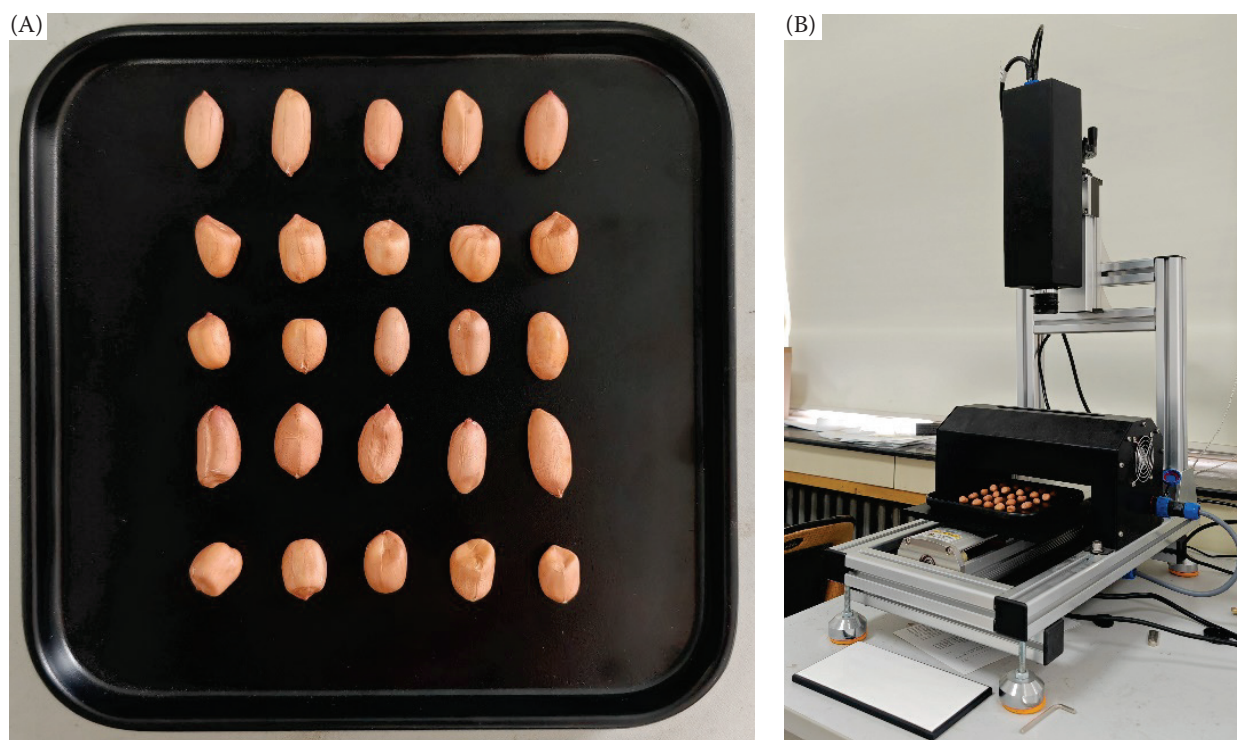


Figure 1. (A) Peanut sample, (B) Hangzhou Caipu Technology Co., Ltd. 'FS-15' series of near-infrared hyperspectral camera

is used to eliminate the remaining small portion of the noise in the image, then the Canny edge detection algorithm (Xuan and Hong 2017) is used to extract the edge of the peanut. And then the hole-filling and corrosion algorithm for final processing of the image is used, the corroded structural unit is taken as a circular struc-

ture with a radius of one. The result of label is obtained by following the above steps.

Extreme learning machine. Extreme learning machine (ELM) is a randomised single-layer feedforward neural network, which is mainly composed of three parts: input layer, hidden layer, and output layer.

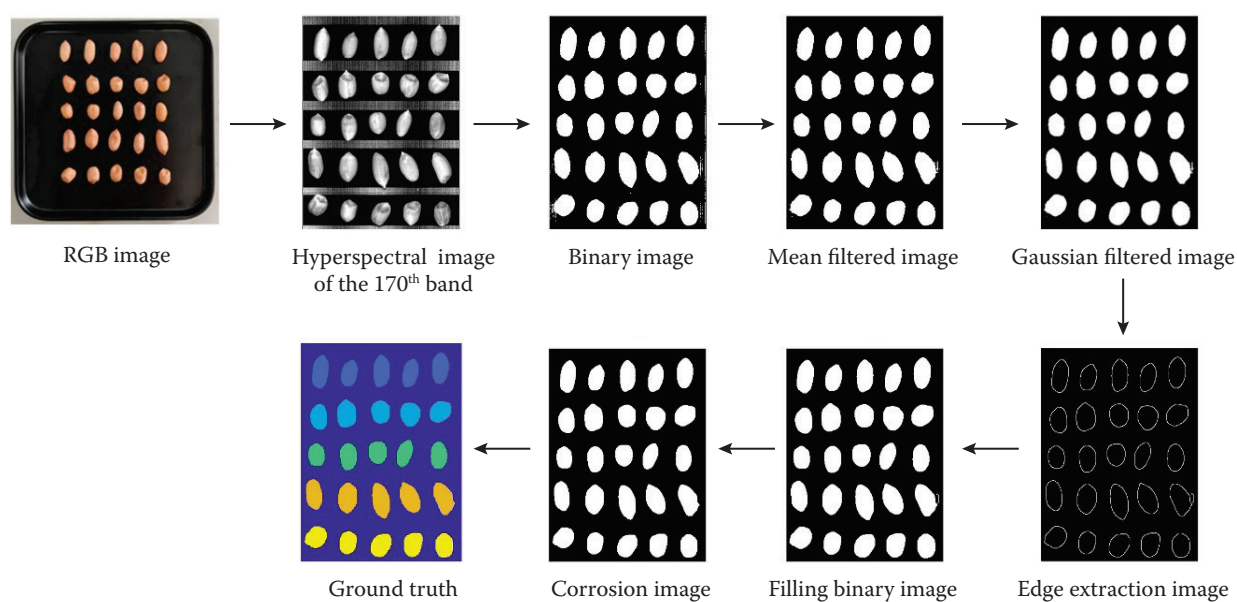


Figure 2. Flowchart of label preparation

<https://doi.org/10.17221/109/2024-CJFS>

The main idea of ELM is to randomly assign the weights between the input layer and the hidden layer, and use quadratic programming problem to calculate the weights of the output layer, which makes the training process of ELM very efficient. The network structure of extreme learning machine is shown in Figure 3.

For hyperspectral data X is a training dataset consisting of n samples:

$$\{(X_j, t_j)\}_{j=1}^N, j=1, \dots, N$$

where: $X_j = [x_{j1}, x_{j2}, \dots, x_{jn}]^T \in R^n$ – j^{th} sample as the n^{th} input attribute of the ELM; $t_j = [t_{j1}, t_{j2}, \dots, t_{jm}]^T \in R^m$ – j^{th} sample as the m^{th} output label.

The number of hidden layer nodes of the ELM is L , then the L hidden layer neurons of the ELM's output function can be written as follows:

$$f_L(X_j) = \sum_{i=1}^L \beta_i g(w_i \times X_j + b_i) = t_j, j=1, \dots, N \quad (2)$$

where: $w_i = [w_{i1}, w_{i2}, \dots, w_{in}]^T \in R^n$, $i = 1, \dots, L$ – input weight vector between the input node and the i^{th} hidden layer node; b_i – bias of the i^{th} hidden layer node; w_i, b_i – randomly generated; $\beta_i = [\beta_{i1}, \beta_{i2}, \dots, \beta_{im}]^T \in R^m$ – output weight vector between the i^{th} hidden layer node and the output node; $g(\cdot)$ – activation function (the activation function is usually a sigmoid function).

The above N equations can be succinctly written as follows (Equations 3–6):

$$H\beta = T \quad (3)$$

$$H = \begin{bmatrix} h(X_1) \\ \vdots \\ h(X_N) \end{bmatrix} = \begin{bmatrix} g(w_1 \times X_1 + b_1) & \dots & g(w_L \times X_1 + b_L) \\ \vdots & \dots & \vdots \\ g(w_1 \times X_N + b_1) & \dots & g(w_L \times X_N + b_L) \end{bmatrix}_{N \times L} \quad (4)$$

$$\beta = \begin{bmatrix} \beta_1^T \\ \vdots \\ \beta_L^T \end{bmatrix}_{L \times m} \quad (5)$$

$$T = \begin{bmatrix} t_1^T \\ \vdots \\ t_N^T \end{bmatrix}_{N \times m} \quad (6)$$

where: $H(X_j) = h(X_j) = [g(w_1 \times X_j + b_1), \dots, g(w_L \times X_j + b_L)]$ – output of the hidden layer neuron for input X_j .

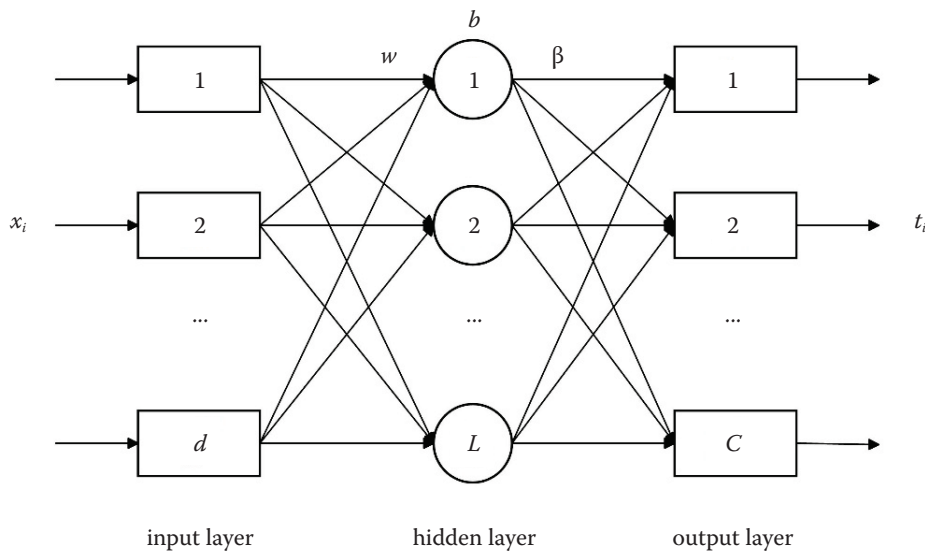


Figure 3. Extreme learning machine network structure diagram

x_i – input sample; w – input weight; b – bias parameter; β – output weight; d – input sample class; L – number of hidden layer nodes; C – output sample class; t_i – output label

It maps the data from the n -dimensional input space to the L -dimensional feature space. The matrix H is the output matrix of the hidden layer and T is the target output matrix. The output weights β are computed by solving a linear least squares problem:

$$\beta = H^+ T \quad (7)$$

where: H^+ – Moore-Penrose generalised inverse of the implicit layer output matrix H .

The Moore-Penrose generalised inverse of H can be computed as $H^+ = H^T(HH^T)^{-1}$. In order to obtain better stability and generality, it is common to add a positive value of $1/\rho$ to each diagonal element of HH^T . The output function of the ELM classifier is thus denoted as:

$$f_L(X_j) = g(X_j)\beta = g(X_j)H^T\left(\frac{1}{\rho} + HH^T\right)^{-1}T \quad (8)$$

Classification of peanut hyperspectral images based on spatial-spectral extreme learning machine

Nonlinear data transformation in high-dimensional feature space increases the probability of linear separability of data in the transformation space. In hyperspectral data, adjacent pixels are generally composed of similar parts, and these similar parts have a high probability of belonging to the same class. It can be known from the ELM algorithm that in the hidden representation process, adjacent pixels in the local window tend to represent the same sample, and the hidden representation features are very close to each other. In order to improve the classification performance of the algorithm, the spatial neighbourhood information in hyperspectral image (HSI) is incorporated into the ELM framework. Currently, several popular spatial-context filters are used in hyperspectral processing, such as propagation filter (Chang and Wang 2015) and bilateral filter (Elad 2002). Therefore, a propagation filter is used in this study to obtain a more accurate hidden representation matrix.

Propagation filtering. Propagation filter is a novel image filtering algorithm, which can not only smooth adjacent image pixels but also preserve image context information such as edges or texture regions without applying explicit spatial kernel function.

For each input data, which is a N -dimensional vector, the ELM maps the data to L -dimensional hidden layer features $H = [h_1, h_2, \dots, h_N]$ through a nonlinear transformation. By using the propagation filter to get spatial information integrated hidden feature representation, which can be described as:

$$\tilde{H} = [PF(h_1), PF(h_2), \dots, PF(h_N)] = [\tilde{h}_1, \tilde{h}_2, \dots, \tilde{h}_N] \quad (9)$$

where: $PF(\)$ – spatial propagation filtering operation; $h_1, h_2, \dots, h_N$ – hidden layer output matrix.

Specifically, the 2D hidden layer matrix $H \in R^{N \times L}$ is reshaped into a 3D cube $T \in R^{M \times W \times L}$, the same as the original hyperspectral data cube, and consider each feature vector as a 'pixel' in the 3D cube, where M is the length of the image, W is the width of the image, and L is the dimension of H . A propagation filter is applied on a 3D cube to extract local contextual spectral spatial features. Let $T(j, k)$ denote the vector T at location (j, k) . The filtered output T produced by the propagation filter is calculated as follows.

$$\hat{T}(j, k) = \frac{1}{Z(j, k)} \sum_{(p, q) \in N(j, k)} w_{(j, k), (p, q)} T(j, k) \quad (10)$$

where: $N(j, k)$ – set of neighbouring pixels centred at (j, k) ; $w_{(j, k), (p, q)}$ – weight of each pixel (p, q) for the centre pixel $T(j, k)$; $Z(j, k)$ – normalisation factor to ensure that the sum of all weights is equal to one.

The algorithm is driven by the idea that, for the pixel (j, k) being related to pixel (p, q) , the intermediate pixels between (j, k) and (p, q) not only need to be photo-metrically related to (j, k) , they are also required to be adjacent-photo-metrically related to their predecessors. As a result, the filter weights are derived by the following definition (Equation 11):

$$w_{(j, k), (p, q)} = w_{(j, k), (p-1, q-1)} \times \exp\left(\frac{-T(p, q) - T(p-1, q-1)}{2\sigma_d^2}\right) \times \exp\left(\frac{-\|T(j, k) - T(p, q)\|^2}{2\sigma_r^2}\right) \quad (11)$$

where: σ_d, σ_r – scale parameters.

<https://doi.org/10.17221/109/2024-CJFS>

As shown in Figure 4, the process of computing the weights $w_{(j,k),(p,q)}$ is demonstrated. After obtaining the propagation-filtered cube data $\hat{T}$, the feature matrix $\hat{H} \in R^{N \times L}$ is obtained by reconstructing the 3D cube $\hat{T} \in R^{M \times W \times L}$ into a 2D matrix, which will be used as the robust hidden feature output of the ELM hidden features.

Classification of hyperspectral images based on spatial-spectral extreme learning machine. In this paper, a hyperspectral image classification method based on spatial spectral limit learning machine is proposed. First, hyperspectral data is put into the input layer of ELM, and the corresponding hidden layer features are learned by input weights. Then the hidden layer features are spatially filtered, the propagation filtering method is used to combine spatial information with spectral information. Finally, the output weights learned from the training set were used to predict the classification results of the test set, to improve the classification accuracy.

Let $X_{\Gamma} \in R^{N_{\Gamma} \times d}$ and $T_{\Gamma} \in R^{N_{\Gamma} \times m}$ be the dataset and labelled set of labelled nodes, respectively, and N_{Γ} be the

number of labelled nodes, and further let $X_u \in R^{N_u \times d}$ be the dataset of unlabelled nodes, where N_u is the number of unlabeled nodes. In order to learn the weights of the output layer, denoted as $\beta \in R^{L \times m}$, the output of the hidden layer is first divided into a labelled part and an unlabeled part, $\hat{H}_{\Gamma}$ and $\hat{H}_u$, respectively. The goal is to assign specific labels to those unlabeled nodes, and for this purpose, the objective function is written as:

$$\hat{H}_{\Gamma} \beta = T_{\Gamma} \quad (12)$$

The expression is reformulated as a regularised ridge regression optimisation problem:

$$\arg \min_{\beta} \ell(\beta) = \arg \min_{\beta} \frac{1}{2} \|\hat{H}_{\Gamma} \beta - T_{\Gamma}\|_F^2 + \frac{\lambda}{2} \|\beta\|_F^2 \quad (13)$$

where: ℓ – loss function; λ – non-negative regularisation factor.

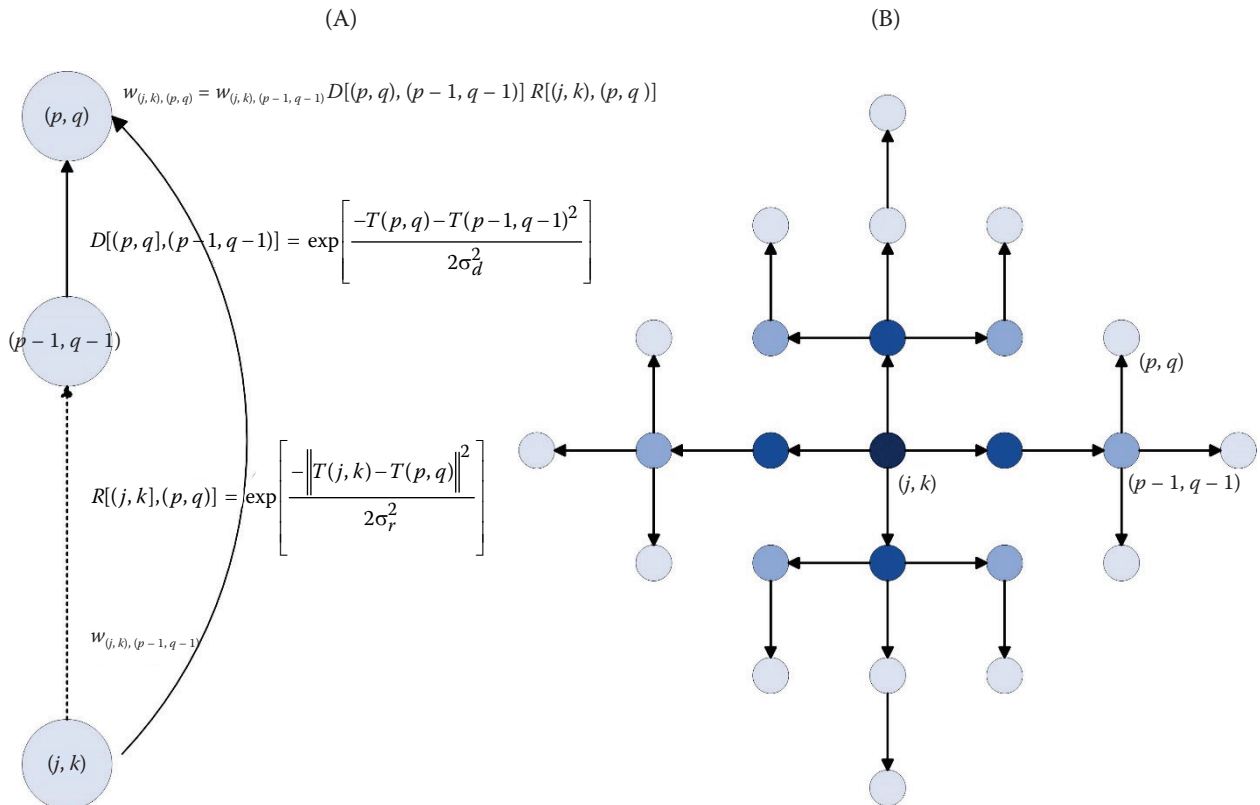


Figure 4. The illustration of the propagation filter: (A) the definition of propagation filtering weight, (B) the pattern of performing 2D propagation filtering with $d = 3$ pixels

(p, q) – edge pixel; (j, k) – centre pixel; w – weight of each pixel; D – weight between the current pixel and the previous pixel; R – weight between the current pixel and the centre pixel; T – vector represented by the centre pixel; σ_d, σ_r – scale parameters

Equation 13 has a closed solution, which can be obtained by calculating the partial derivatives of ℓ with respect to β . The partial derivative can be written as:

$$\frac{\partial \ell}{\partial \beta} = \hat{H}_r^T H_r \beta + \lambda \beta - \hat{H}_r^T T_r \quad (14)$$

Setting Equation 14 to 0 results in the following solution:

$$\beta = (\hat{H}_r^T H_r + \lambda I_L)^{-1} \hat{H}_r^T T_r \quad (15)$$

As a result, the labels of the unlabeled nodes can be determined:

$$T_u = \hat{H}_u \beta \quad (16)$$

RESULTS AND DISCUSSION

Experimental setup

Five visually clean, intact, uniformly sized, similarly shaped, and closely coloured peanuts were selected from each variety as experimental samples. The samples were arranged on a 17.5×17.5 cm black tray, and data acquisition was performed using a spectral camera. The size of dataset is 601×320 with 256 bands. After removing the edge, a subset is adopted, which contains of 434×320 pixels and 256 bands. The dataset contains 37 905 samples and is classified into five categories. The spectral curve is shown in Figure 5.

In order to verify the classification performance of the method, several different classification methods are compared: support vector machine (SVM) (Chang and Lin 2011), sparse multinomial logistic regression (SMLR) (Krishnapuram et al. 2005), sparse multinomial logistic regression with multi-level logic spatial priors (SMLR-MLL) (Li et al. 2010), and sparse multinomial logistic regression with spatially adaptive total

variation (SMLR-SpATV) (Sun et al. 2014), extreme learning machine (ELM) (Huang et al. 2006, 2011, 2015), bilateral filtering-extreme learning machine (BF-ELM). All classification algorithms are implemented using MATLAB (version R2022b) on an Intel i7-13700HX 2.1 GHz CPU, RTX 4080 12 GB GPU, and 32 GB RAM.

The parameters of all the comparison algorithms refer to the values provided in the original text. For the proposed method, the scale parameters σ_r and d_r in the propagation filtering and bilateral filtering methods are uniformly set to 0.8 and 2.

In the experimental results, overall accuracy (OA), average accuracy (AA), category accuracy (CA), and kappa coefficient are employed to evaluate the performance of different classification methods. Table 1 shows the classification performance of various algorithms.

Parameter analysis

Number of training samples. 1% to 10% labelled samples are used in this study to train our model, and comparisons are made with other algorithms to explore the classification performance with different numbers of training samples. As can be seen from Figure 6A–C, with the increase of the number of training samples, each algorithm can obtain more comprehensive information in the process of learning features, so the OA, AA, and kappa of each algorithm are also improved. The method proposed in this paper has the highest classification accuracy and the best effect, which is superior to other algorithms. In this experiment, the number of training samples is 3 411.

Number of neurons in the hidden layer. In Figure 7A, the influence of the number of neurons in the hidden layer of ELM, BF-ELM, and PF-ELM on the classification results is shown. The average accuracy for the number of neurons in the hidden layer is plotted from 100 to 1 000 in intervals of 100. It can

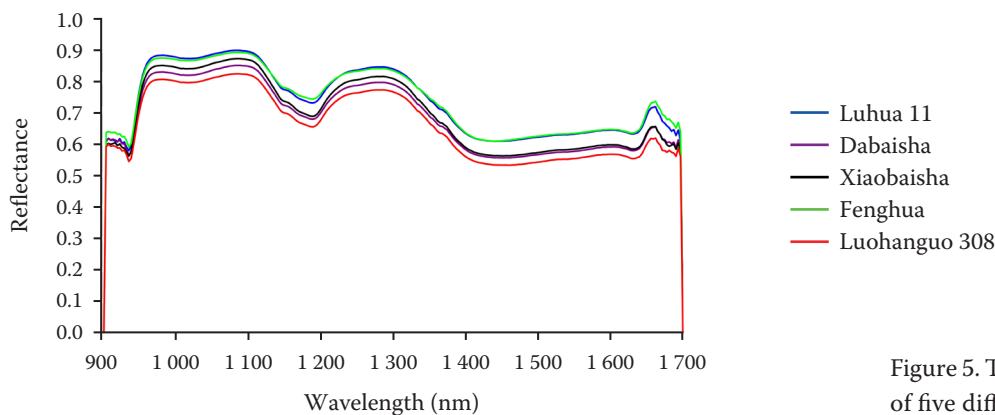


Figure 5. The spectral values curves of five different peanut varieties

<https://doi.org/10.17221/109/2024-CJFS>

Table 1. Classification accuracy of each algorithm

Class	SVM	ELM	BF-ELM	PF-ELM	SMLR	SMLR-MLL	SMLR-SpATV
1 Luhua 11	87.44	54.59	91.98	98.30	62.96	82.96	94.00
2 Dabaisha	83.87	43.37	91.66	99.12	52.02	70.68	96.70
3 Xiaobaisha	87.65	47.85	93.41	98.68	59.85	78.51	94.40
4 Fenghua	90.25	70.78	93.69	98.00	74.73	86.41	100.00
5 Luohanguo 308	86.27	56.73	95.16	97.57	59.86	81.76	97.38
OA (%)	87.23	55.49	93.16	98.32	62.52	80.37	96.66
AA (%)	87.09	54.67	93.18	98.33	61.88	80.07	96.49
Kappa	0.840	0.442	0.914	0.979	0.530	0.754	0.958

OA – overall accuracy; AA – average accuracy; SVM – support vector machine; ELM – extreme learning machine; BF-ELM – bilateral filtering-extreme learning machine; PF-ELM – propagation filtering-extreme learning machine; SMLR – sparse multinomial logistic regression; SMLR-MLL – sparse multinomial logistic regression with multi-level logic spatial priors; SMLR-SpATV – sparse multinomial logistic regression with spatially adaptive total variation

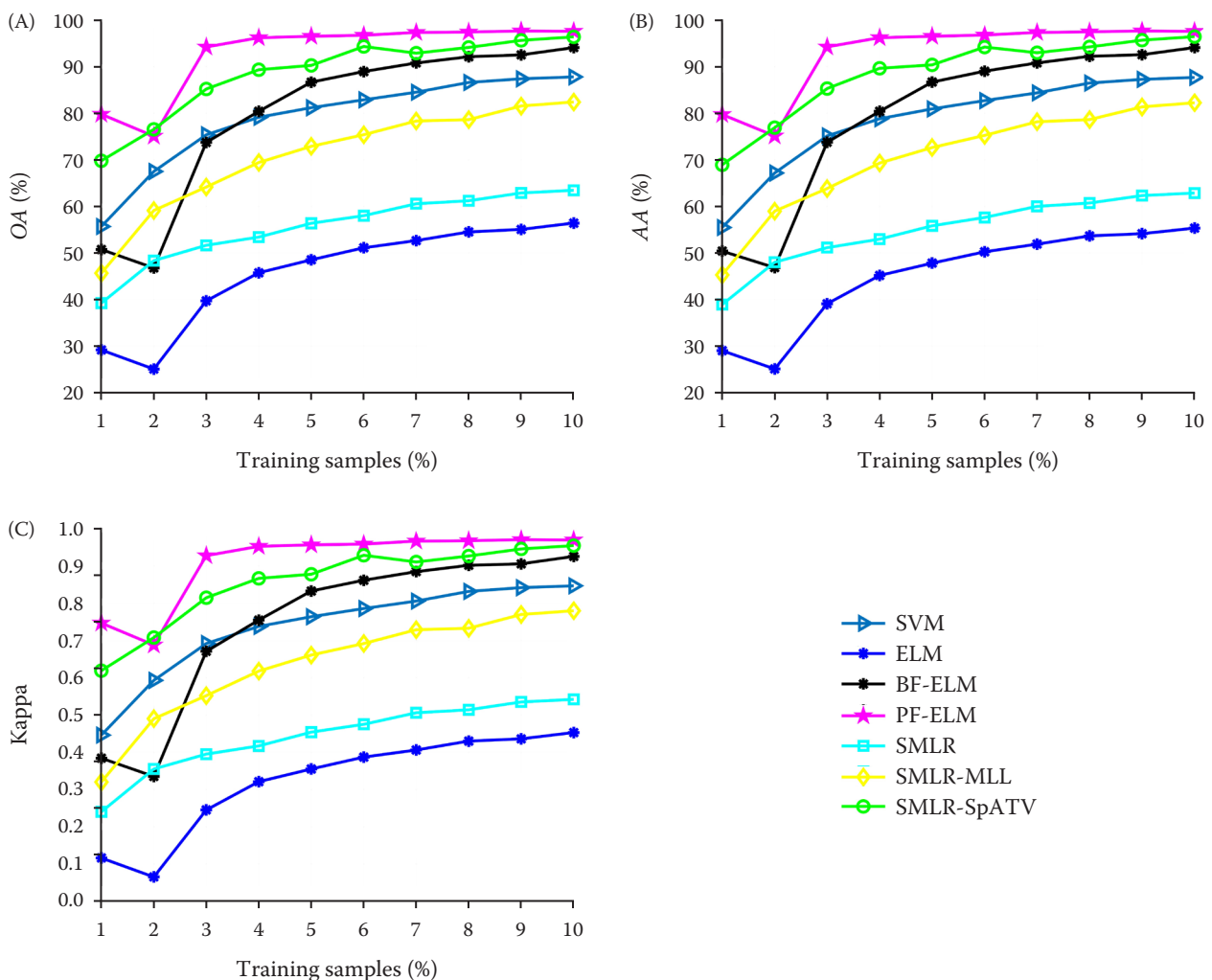


Figure 6. The influence of the number of training samples on the classification results

OA – overall accuracy; AA – average accuracy; SVM – support vector machine; ELM – extreme learning machine; BF-ELM – bilateral filtering-extreme learning machine; PF-ELM – propagation filtering-extreme learning machine; SMLR – sparse multinomial logistic regression; SMLR-MLL – sparse multinomial logistic regression with multi-level logic spatial priors; SMLR-SpATV – sparse multinomial logistic regression with spatially adaptive total variation

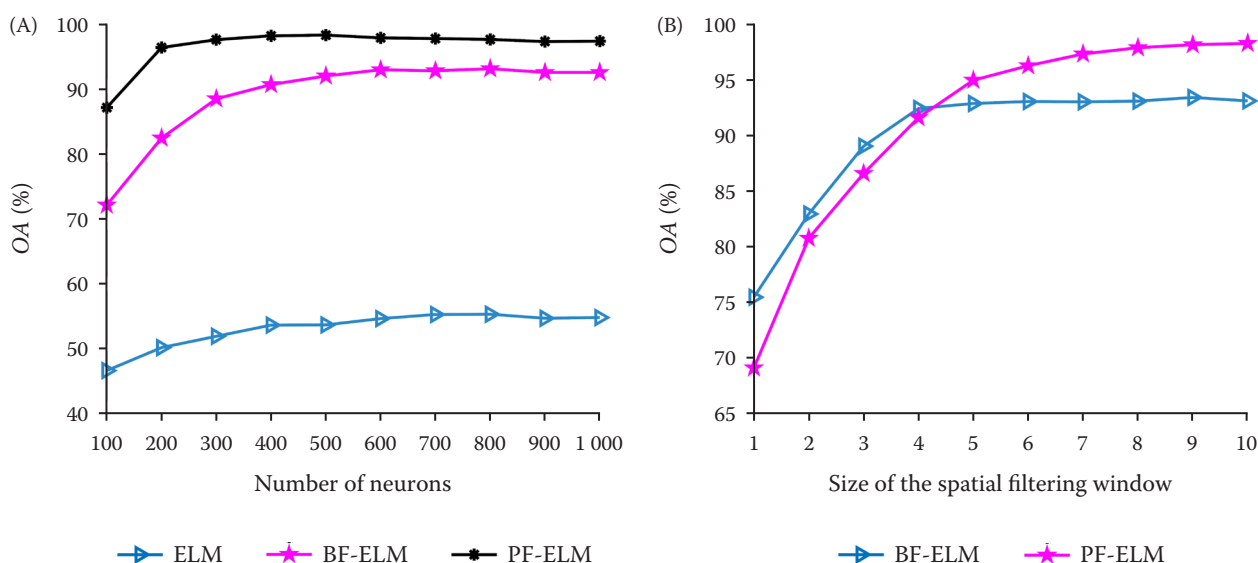


Figure 7. The influence of the number of neurons in the hidden layer and the window size of spatial filtering on the classification results

OA – overall accuracy; ELM – extreme learning machine; BF-ELM – bilateral filtering-extreme learning machine; PF-ELM – propagation filtering-extreme learning machine

be seen from the figure that when the number of neurons is 500, the classification accuracy of PF-ELM is the highest. As the number of neurons increases, the classification accuracy of PF-ELM is generally better than that of ELM and BF-ELM. It should be noted that when the number of hidden layer neurons is too large, the classification accuracy may be reduced. This is because increasing the number of neurons in the hidden layer also increases the risk of overfitting about training data. Therefore, in the experiment, the value of hidden layer nodes in ELM and BF-ELM algorithms is 800, and the value of hidden layer nodes in PF-ELM algorithm is 500.

Window size for spatial filtering. The effect of the proposed SS-ELM on the classification accuracy of spatial filtering with different window sizes will be discussed. It can be seen from Figure 7B that the classification accuracy of PF-ELM improves as the window size increases, but slowly tends to balance to some extent. For the BF-ELM algorithm, as the window size starts to become larger, the classification accuracy basically tends to balance. Therefore, in this experiment, the spatial filtering window size of the two algorithms is 10.

Analysis of results

In the experiments, the classification accuracy of the proposed method was assessed by comparing it with other classification methods. Table 1 presents the OA, AA, kappa coefficients, and accuracy for each category for all algorithms. Additionally, Figure 8B–H illustrates

the visual performance of the classification results for all algorithms.

i) As shown in Table 1, ELM gets the lowest classification accuracy, with an OA of only 55.49%. Following closely is SMLR, which also shows a relatively low OA of 62.52%, falling below the 80% for classification accuracy. The OA of SVM is only 87.23%.

ii) The performance of SMLR-MLL and SMLR-SPATV is enhanced by incorporating spatial information from hyperspectral data. These two methods demonstrate a significant improvement in classification accuracy. Specifically, the OA of SMLR-MLL achieves 80.37%, which is 17.85% higher than that of SMLR. Additionally, the AA increases by 18.19% and kappa increases by 0.224 with the use of SMLR-MLL. On the other hand, the OA of SMLR-SPATV reaches an impressive 96.66%. In comparison to SMLR, the OA of SMLR-SPATV shows a remarkable increase of 34.14%, while AA increases by 34.61% and kappa increases by 0.428.

iii) The classification accuracy of BF-ELM and PF-ELM is improved by adding spatial information. The OA of BF-ELM gets 93.16%, which is 37.67% higher than that of ELM, AA increases 38.51%, and kappa increases 0.472. The OA of PF-ELM gets 98.32%, which is 42.83%, 43.66% and 0.537 higher than the OA, AA, and kappa of ELM. For BF-ELM and PF-ELM, it can be seen that PF-ELM is better than BF-ELM, OA is increased by 5.16%, AA is increased by 5.15%, and kappa is increased by 0.065.

<https://doi.org/10.17221/109/2024-CJFS>

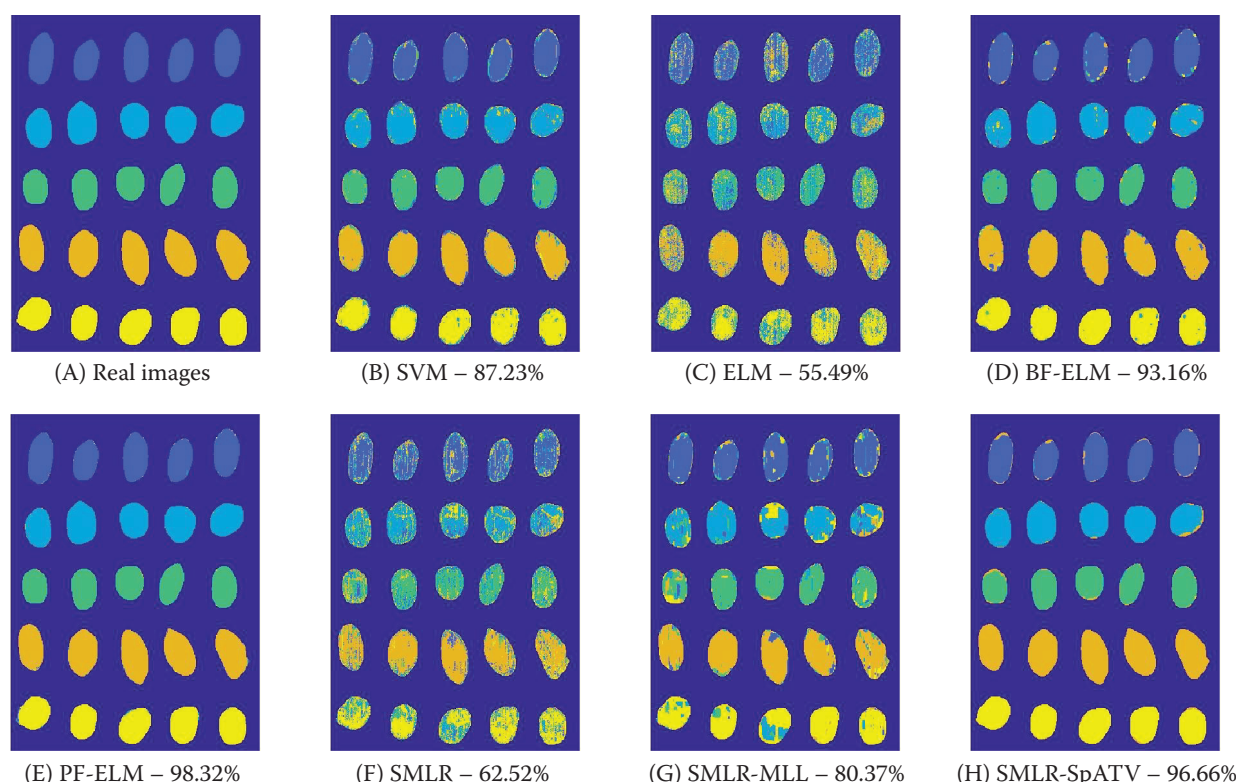


Figure 8. Classification performance and visualisation of each algorithm: (A) real images, (B) SVM, (C) ELM, (D) BF-ELM, (E) PF-ELM, (F) SMLR, (G) SMLR-MLL, (H) SMLR-SpATV

SVM – support vector machine; ELM – extreme learning machine; BF-ELM – bilateral filtering-extreme learning machine; PF-ELM – propagation filtering-extreme learning machine; SMLR – sparse multinomial logistic regression; SMLR-MLL – sparse multinomial logistic regression with multi-level logic spatial priors; SMLR-SpATV – sparse multinomial logistic regression with spatially adaptive total variation

iv) The classification performance of PF-ELM is better than that of SVM, OA is increased by 11.09%, AA is increased by 11.24%, and kappa is increased by 0.139.

v) Compared to SMLR-MLL and SMLR-SpATV, PF-ELM demonstrates superior capability in capturing context information between pixels in the image. At the same time, the iterative calculation process is also eliminated, thereby saving time and reducing computational workload. Furthermore, PF-ELM exhibits a significant improvement in classification accuracy. In comparison to SMLR-MLL, PF-ELM shows an increase of 17.95% in OA, 18.26% in AA, and 0.225 in kappa. When compared to SMLR-SpATV, PF-ELM achieves a 1.66% increase in OA, a 1.84% increase in AA, and a 0.021 increase in kappa.

CONCLUSION

In this paper, a hyperspectral image classification method based on spatial-spectral extreme learning machine (SS-ELM) is proposed by the authors to classify

peanut varieties quickly and accurately. The method inherits all the advantages from ELM, a local spectral-spatial context integration and reshaping mechanism is incorporated into the hidden layer feature representation by using a context-aware propagation filtering procedure. The experimental results show that the accuracy of the improved ELM model on five varieties of peanuts dataset (Luhua 11, Dabaisha, Xiaobaisha, Fenghua, and Luohanguo 308) was 98.32%, which was higher than other classic models, proving the feasibility of hyperspectral imaging and ELM in peanut variety identification and classification.

REFERENCES

- Carrin M.E., Carelli A.A. (2010): Peanut oil: Compositional data. *European Journal of Lipid Science and Technology*, 112: 697–707.
- Chang C.C., Lin C.J. (2011): LIBSVM: A library for support vector machines. *ACM Transactions on Intelligent Systems and Technology*, 2: 1–27.

- Chang J.H.R., Wang Y.C.F. (2015): Propagated image filtering. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Boston, USA, June 7–12, 2015: 10–18.
- Deng G., Cahill L.W. (1993): An adaptive Gaussian filter for noise reduction and edge detection. In: 1993 IEEE Conference Record Nuclear Science Symposium and Medical Imaging Conference, San Francisco, USA, Oct 31–Nov 6, 1993: 1615–1619.
- Elad M. (2002): On the origin of the bilateral filter and ways to improve it. *IEEE Transactions on Image Processing*, 11: 1141–1151.
- Elmasry G., Kamruzzaman M., Sun D.W., Allen P. (2012): Principles and applications of hyperspectral imaging in quality evaluation of agro-food products: A review. *Critical Reviews in Food Science and Nutrition*, 52: 999–1023.
- Fabiyi S.D., Vu H., Tachtatzis C., Murray P., Harle D., Dao T.K., Andonovic I., Ren J., Marshall S. (2020): Varietal classification of rice seeds using RGB and hyperspectral images. *IEEE Access*, 8: 22493–22505.
- Gupta G. (2011): Algorithm for image processing using improved median filter and comparison of mean, median and improved median filter. *International Journal of Soft Computing and Engineering*, 1: 304–311.
- Huan L., Yaqian W., Xiaoming W., Dong A., Yaoguang W., Laixin L., Xing C., Yanlu Y. (2019): Study on detection method of wheat unsound kernel based on near-infrared hyperspectral imaging technology. *Spectroscopy and Spectral Analysis*, 39: 223–229.
- Huang G.B., Zhu Q.Y., Siew C.K. (2006): Extreme learning machine: Theory and applications. *Neurocomputing*, 70: 489–501.
- Huang G.B., Zhou H., Ding X., Zhang R. (2011): Extreme learning machine for regression and multiclass classification. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 42: 513–529.
- Huang G., Huang G.B., Song S., You K. (2015): Trends in extreme learning machines: A review. *Neural Networks*, 61: 32–48.
- Jiang H., Liu L., Chen Q. (2022): Rapid determination of acidity index of peanuts by near-infrared spectroscopy technology: Comparing the performance of different near-infrared spectral models. *Infrared Physics and Technology*, 125: 104308.
- Krishnapuram B., Carin L., Figueiredo M.A., Hartemink A.J. (2005): Sparse multinomial logistic regression: Fast algorithms and generalization bounds. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 27: 957–968.
- Li J., Bioucas-Dias J.M., Plaza A. (2010): Exploiting spatial information in semi-supervised hyperspectral image segmentation. In: 2010 2nd Workshop on Hyperspectral Image and Signal Processing: Evolution in Remote Sensing, Reykjavik, Iceland, June 14–16, 2010: 1–4.
- Liu Z., Jiang J., Qiao X., Qi X., Pan Y., Pan X. (2020): Using convolution neural network and hyperspectral image to identify moldy peanut kernels. *LWT – Food Science and Technology*, 132: 109815.
- Liu Q., Wang Z., Long Y., Zhang C., Fan S., Huang W. (2022): Variety classification of coated maize seeds based on Raman hyperspectral imaging. *Spectrochimica Acta Part A: Molecular and Biomolecular Spectroscopy*, 270: 120772.
- Mohi-Alden K., Omid M., Firouz M.S., Nasiri A. (2023): A machine vision-intelligent modelling based technique for in-line bell pepper sorting. *Information Processing in Agriculture*, 10: 491–503.
- Qu J.H., Liu D., Cheng J.H., Sun D.W., Ma J., Pu H., Zeng X.A. (2015): Applications of near-infrared spectroscopy in food safety evaluation and control: A review of recent research advances. *Critical Reviews in Food Science and Nutrition*, 55: 1939–1954.
- Sun L., Wu Z., Liu J., Xiao L., Wei Z. (2014): Supervised spectral-spatial hyperspectral image classification with weighted Markov random fields. *IEEE Transactions on Geoscience and Remote Sensing*, 53: 1490–1503.
- Wang L., Liu D., Pu H., Sun D.W., Gao W., Xiong Z. (2015): Use of hyperspectral imaging to discriminate the variety and quality of rice. *Food Analytical Methods*, 8: 515–523.
- Wright H. (1980): Commercial hybrid seed production. In: Fehr W.R., Hadley H.H. (eds): *Hybridization of Crop Plants*. Madison, American Society of Agronomy and Crop Science Society of America: 161–176.
- Wu Q., Xu L., Zou Z., Wang J., Zeng Q., Wang Q., Zhou M. (2022): Rapid nondestructive detection of peanut varieties and peanut mildew based on hyperspectral imaging and stacked machine learning models. *Frontiers in Plant Science*, 13: 1047479.
- Xuan L., Hong Z. (2017): An improved canny edge detection algorithm. In: 2017 8th IEEE International Conference on Software Engineering and Service Science (ICSESS), Beijing, China, Nov 24–26, 2017: 275–278.
- Yuan D., Jiang J., Qi X., Xie Z., Zhang G. (2020): Selecting key wavelengths of hyperspectral image for nondestructive classification of moldy peanuts using ensemble classifier. *Infrared Physics and Technology*, 111: 103518.
- Zhang L., Wang Y., Wei Y., An D. (2022): Near-infrared hyperspectral imaging technology combined with deep convolutional generative adversarial network to predict oil content of single maize kernel. *Food Chemistry*, 370: 131047.
- Zou Z., Wang L., Chen J., Long T., Wu Q., Zhou M. (2022): Research on peanut variety classification based on hyperspectral image. *Food Science and Technology*, 42: e18522.

Received: June 7, 2024

Accepted: December 16, 2024

Published online: January 29, 2025