

An Improved Face Recognition Method Using Optimized Subspace Decomposition and RBF Neural Networks

Saleem Pasha Choudhary R¹, Riyazoddin S², Mohammed Ahmed³, Md Ameenuddin⁴

^{1,2,4}Theem College of Engineering, Boisar

³M.H. Saboo Siddik College of Engineering, Byculla, Mumbai

Abstract:- Face biometrics is one of the most reliable tools used in applications such as personal identification and authentication, surveillance systems, and human–computer interaction. However, we usually observe the significant influence on recognition performance of these applications due to poor quality inputs affected by illumination, variations in facial pose, and overfitting challenges because of large dimensionality of image data. In order to mitigate the performance challenge, this study presents a optimized intelligent subspace decomposition-based framework to enhance the face recognition accuracy. In this framework integration of PCA, LDA, SVD, and FrSVD is used for effective feature representation and dimensionality reduction. In addition to this a weighted feature fusion mechanism is employed to integrate information obtained from the different feature extraction techniques to enhance the quality of discriminative characteristics of the resulting feature representation. The fused feature vector is subsequently classified using a Radial Basis Function (RBF) neural network. For experimentation, Yale B face data base is used to evaluate the performance of this improved framework thereby attaining the improved 96.5% of recognition accuracy compared to conventional PCA–LDA-based methods, specifically, under variations in illumination and facial pose.

Keywords: Face Recognition, Principal Component Analysis, Linear Discriminant Analysis, Singular Value Decomposition, Fractional Singular Value Decomposition, Feature Fusion, RBF Neural Network.

Introduction

In recent past, after covid-19 pandemic the face biometric has attracted the considerable attentions because of its increasing use in real-time applications [1,4,6] which includes biometric authentication, access control, surveillance, and human–computer interaction. Even though, we observe a substantial progress has been achieved in these applications through machine learning and deep learning techniques but reliable face recognition remains a challenge in real time practical environments. It is mostly due to the factors such as noisy observations, high-dimensional image representations, changes in facial appearance, pose variations, occlusions, and environmental conditions which adversely affects recognition performance. In recent [2,3, 20-22] there is considerable performance degradation when facial images are acquired under difficult illumination, occlusion, resolution, and appearance conditions. In particular, recognition under low-light conditions continues to pose difficulties even for contemporary learning-based systems there by highlighting the need for feature representations that are more resistant to illumination changes [22].

In face classification task, the effectiveness of a classifier is commonly evaluated according input size within a feature space and the ability to establish decision boundaries that provide adequate separation among different classes. According to [7,8] the classification accuracy is good if the samples belong to different classes with limited varying features

In face recognition, however [20], the facial inputs of same person may differ substantially due to changes in illumination, pose, and facial expression. Subsequently, the features classification may of different individuals may overlap giving the very poor performance of the system. This varying condition has motivated to find robust feature extraction, dimensionality-reduction, and classification techniques

thereby improving the recognition accuracy without imposing excessive computational requirements.

A conventional face recognition system can broadly be divided into two principal phases. The first phase involves detecting and localizing the face, followed by suitable preprocessing and extraction of representative features from the detected facial region [8,10]. The second phase resulting feature vector is deciding appropriate classifier with good classification algorithm [6,11]. However, the reliability of this phase significantly gets influenced by factors such as facial position, scale, orientation, illumination, and expression. Thus, the effectiveness of the complete recognition system is strongly dependent on both the discriminative quality of the extracted representation and the capability of the selected classifier. So, the recent research [20-22] for recognition under varying conditions has continued to emphasize the importance of representations that preserve identity-related characteristics while suppressing variations that are unrelated to identity.

Earlier Turk and Pentland [1] in face recognition has introduced Principal Component Analysis (PCA) among the most established statistical approach employed for feature extraction and dimensionality reduction. In this technique high-dimensional facial image data is projected onto a lower-dimensional subspace to form an eigenvector corresponding to the largest eigenvalues of the covariance matrix. Thus, PCA provides a compact representation while reducing the dimensionality of the original image data by retaining the directions associated with the greatest variance which in-turn reduces the computational burden.

However, PCA is an unsupervised dimensionality-reduction technique and does not explicitly incorporate class-label information during the construction of the feature space. As a result, variations that contribute significantly to overall data variance may not necessarily improve class discrimination. This limitation can become more pronounced when facial images contain substantial natural or environmental variations [20].

Linear Discriminant Analysis (LDA) [10] can be employed to address the class-discrimination limitation of PCA. In contrast to PCA, LDA uses class information to seek a feature space in which the separation between classes is increased while the variation within each class is reduced. Applying LDA to a PCA-reduced representation can therefore produce facial features with stronger discriminatory properties. For this reason, PCA combined with LDA has continued to serve as an important baseline in appearance-based face recognition. Additional representation of facial image information can be obtained through Singular Value Decomposition (SVD) [6,10,15]. SVD decomposes an image matrix into a set of orthogonal components characterized by their corresponding singular values, thereby providing a compact representation of the underlying image information. The decomposition can be useful for identifying and retaining significant components while reducing the influence of less informative information.

Liu, et al. [15] has suggested a Fractional Singular Value Decomposition Representation (FSVDR) as an extension of the conventional SVD-based representation through the incorporation of a fractional-order parameter that controls the contribution of singular-value components [23]. According to the study of FrSVD, most of the times the dominant singular features may be influenced due the varying facial features because of variations such as illumination, expression, and occlusion. The use of fractional component inflates the lower dimensional poor-quality features and deflates the higher dimensional features there by getting the uniform quality features that can help to reduce the adverse effects of such variations. In [23] usefulness of FSVDR in conjunction with dimensionality-reduction techniques such as PCA and LDA is modified to produce the acceptable recognition accuracy compared to the conventional PCA- and LDA-based representations particularly when facial images undergo considerable appearance changes. The findings [20-26] have ensured that a appropriate facial feature representation gives the improved recognition accuracy there by maintaining the original identity. These methods mostly include the residual feature decomposition and multi-task learning framework for variation-invariant face recognition by reducing the influence of illumination and other undesirable variations. Some of these methods [20-26] have suggested the integration of multiple complementary feature representations obtained from different descriptors to provide single combined common representation, potentially providing spatial, structural, spectral, and

discriminative information that may not be adequately captured by a single feature extraction method. These research work has significance of combining complementary feature sets for face recognition in the presence of illumination, pose, expression, and noise variations [24]. More recent fusion-based investigations [25] have studied and analysed the integration of complementary facial representations to improve recognition robustness.

In the proposed improved framework, PCA is used to obtain the appearance information, LDA for class-discriminative characteristics LDA, and spectral information using SVD and FrSVD to produce the fused unified feature representation. This will avoid dependency of any single feature representation of complex face features.

This complex feature vector is further classified using the deep learning techniques since neural based classifiers have found to be achieving reliable recognition accuracy. So, in proposed frame work Radial Basis Function (RBF) neural networks are used as it possesses compact network structure [1,4,5]. Also, RBF network uses a localized radial basis activation function to transform the input feature space and can subsequently establish nonlinear decision boundaries between different classes [26,27] retaining the original quality of fused features. In proposed framework a predefined weighting coefficients are employed to control the contribution of each feature representation. The weighted fusion process is intended to balance the different sources of information and produce a unified feature vector suitable for subsequent classification in the resulting feature space affected due to the presence of illumination and other facial appearance variations.

The performance of the proposed framework is evaluated using the Yale B face database [19], containing facial images acquired under significantly different illumination conditions and achieves a recognition accuracy of 96.5%. The experimental results indicate that integrating PCA, LDA, SVD, and FrSVD through weighted feature fusion can provide improved recognition performance compared with the conventional PCA–LDA approach. The findings further suggest that the proposed framework can serve as a relatively simple and computationally efficient approach for generating discriminative facial representations when illumination-related variations are present.

The major contributions of this work include:

1. A improved hybrid framework integrating PCA, LDA, SVD, and FrSVD is developed for facial feature extraction and dimensionality reduction.
2. A simple weighted feature fusion optimization is carried out to integrate complementary information generated from different feature representation techniques.
3. An RBF neural network classifier is utilized to classify the resulting fused feature vectors by establishing nonlinear decision boundaries.
4. The proposed framework is experimentally evaluated by using the Yale B face database and obtained an improved recognition accuracy of 96.5% under the evaluated experimental conditions.

The remainder of this paper is structured as follows. Section 2 reviews existing research related to face recognition, dimensionality reduction, feature extraction, feature fusion, and RBF neural-network-based classification. Section 3 details the proposed methodology, including image preprocessing, PCA, LDA, SVD, FrSVD, weighted feature fusion, and RBF-based classification. Section 4 describes the experimental configuration, database, evaluation measures, and implementation procedure also analyses the experimental results and compares the performance of the proposed framework with existing face recognition approaches. Finally, Section 5 discusses conclusion, limitations and potential directions for future research.

2. RELATED WORK

Due growing importance of face recognition in smart applications used for biometric authentication, access control, surveillance, identity verification, and human–computer interaction, face recognition has been extensively studied as a pattern-recognition problem. Most conventional face recognition systems first reduce the dimensionality of the high-dimensional facial image data into a lower-dimensional representation that captures

the most relevant information and then perform the classification. In this context, a number of different techniques such as Principal Component Analysis (PCA), Linear Discriminant Analysis (LDA), Singular Value Decomposition (SVD), fractional-order representations, local feature descriptors, and neural-network-based classifiers have been explored to improve the recognition accuracy, robustness and computational efficiency [1, 3-6].

2.1. PCA-Based Face Recognition

Turk and Pentland [1] have first time used the Principal Component Analysis a well-established statistical approach for dimensionality reduction of facial image data and constructing appearance-based representations. They have introduced the Eigenface based method, where facial images are represented in a reduced eigenspace obtained from the covariance characteristics of the training samples [1]. By retaining the eigenvectors corresponding to the largest eigenvalues, PCA produces a compact representation that preserves the dominant variance present in the original facial image space.

An important limitation of PCA is that it is an unsupervised transformation and therefore does not explicitly exploit information about the identity classes. The directions associated with maximum variance may consequently correspond to variations that are not useful for distinguishing between subjects. This issue becomes more significant when changes in illumination, facial pose, expression, or environmental conditions introduce substantial variations into the image data.

To overcome the limitations of linear subspace representations, nonlinear variants of PCA have also been explored. Kim et al. developed a Kernel PCA-based face recognition method in which facial observations are implicitly transformed into a nonlinear feature space before principal component analysis is applied [9]. Such nonlinear mappings can represent facial patterns whose distributions cannot be adequately described using a conventional linear subspace. However, kernel-based methods may need extra computation and careful tuning of kernel-related parameters. At the same time improving the efficiency and reliability of facial feature representation is also an important research objective. Zhang and Yan [2] proposed a fast face recognition method using image-gradient compensation to improve facial feature representation and recognition efficiency. In general, previous studies have shown that changes in illumination, pose, expression and image quality can greatly affect the reliability of extracted facial features [7]. These limitations motivate the use of other discriminative feature extraction mechanisms after the initial dimensionality-reduction stage.

2.2. LDA and Discriminative Feature Extraction

Linear Discriminant Analysis provides a supervised approach to feature extraction by incorporating class information during the construction of the feature space. Belhumeur et al. introduced the Fisherfaces method, which employs Fisher's discriminant criterion to increase the separation between different classes while reducing variations among samples belonging to the same class [3]. Consequently, unlike PCA, LDA is explicitly designed to generate features with stronger class-discriminative properties.

LDA has also been combined with local feature descriptors to improve the representation of facial structures. For instance, Gabor-based features integrated with LDA have been investigated for capturing local spatial and orientation characteristics while maintaining class discrimination [10], [11]. Such representations are useful because identity-related information can be distributed over localized regions of the face and may not be adequately captured by global appearance-based representations alone.

One of the principal difficulties associated with conventional LDA is the small-sample-size problem. Facial images commonly contain a very large number of dimensions compared with the number of available training samples. Under these circumstances, the within-class scatter matrix may become singular or poorly conditioned, which can make direct application of LDA numerically unstable. A commonly adopted solution is therefore to perform dimensionality reduction before applying LDA. The PCA-LDA sequence first reduces the dimensionality of the feature space and then performs supervised discriminant analysis, making it an important conventional approach for appearance-based face recognition [3], [8], [9].

Although PCA-LDA can generate an effective discriminative representation, the resulting feature space is mainly determined by statistical variance and class-separation information. Other structural and spectral representations may therefore contain complementary characteristics that are not completely represented by PCA-LDA.

2.3. SVD-Based Face Representation

Singular Value Decomposition provides another mathematical means of obtaining compact representations from facial images and feature matrices. Through SVD, an input matrix is expressed in terms of orthogonal basis components and their associated singular values. This decomposition makes it possible to represent the dominant structural characteristics of the input using a smaller number of components. Several studies have examined SVD-based representations either independently or together with discriminant-analysis techniques. Deswal et al. investigated an SVD-LDA face recognition framework in which a singular-value-based representation was combined with LDA to enhance the separation of facial patterns [12]. Lu and Zhao also studied dominant SVD representations for face recognition and demonstrated the usefulness of the major singular components for obtaining compact facial representations [18].

An important property of SVD is that it can produce low rank approximations that still preserve the major structural information of the original data. The magnitude of each singular value indicates the relative contribution of its corresponding component, and when a reduced representation is in order, components with relatively low significance can be discarded. However, illumination and pose conditions can significantly alter the global structure of facial images and impact the main singular components.

Thus, conventional SVD can be seen as an additional feature representation rather than an independent solution for robust face recognition in varying conditions. The combination of SVD with discriminative methods and fractional order representations provides an opportunity to exploit information from different feature spaces.

2.4. Fractional SVD-Based Representation

Liu et al [15] proposed a Fractional Singular Value Decomposition Representation (FSVDR) method for face recognition by elevating the conventional SVD representation. These methods had attracted attention of the research community from signal and image processing domain since it offer additional flexibility in controlling the contribution of individual components within a representation. According to the studies the conventional SVD uses the singular values as features obtained from the matrix decompositions, However, the least significant singular value may get suppressed due to dominant values there by affecting the recognition accuracy. On the other hand, the fractional parameter of FrSVD method provides the uniform features vector of the same facial information. It also mitigates the influence of sensitive facial changes caused due to illumination, expression, and occlusion [15]. Thus, Conventional SVD primarily emphasizes dominant structural components, whereas the fractional transformation adjusts the relative contribution of the corresponding spectral components.

2.5. Feature Fusion for Face Recognition

Feature fusion is an effective strategy for combining complementary information generated by multiple feature extraction techniques. Yang et al. examined parallel and serial feature-fusion approaches and showed that the manner in which different feature representations are integrated can affect recognition performance [22]. The underlying principle of feature fusion is that different descriptors can characterize different properties of the same facial image. Combining them may therefore produce a richer representation than that obtained from an individual descriptor.

Within the proposed framework, PCA and LDA generate statistical and discriminative descriptions of the facial data, respectively, whereas SVD and FrSVD contribute complementary structural and fractional-order spectral information. Since these transformations characterize the facial image from different perspectives, their combination can potentially compensate for limitations associated with an individual representation.

Feature-fusion techniques have also been investigated in face recognition systems operating under challenging imaging conditions [24]–[27]. Although complex fusion architectures can generate highly optimized representations, they may require additional parameters and computational resources. Weighted feature fusion provides a comparatively straightforward alternative by explicitly determining the contribution of each feature representation through predefined coefficients. In the proposed framework, the normalized feature representations are integrated using weighting coefficients that satisfy

$$w_1 + w_2 + w_3 = 1 \quad (2.1)$$

The resulting feature representation combines the discriminative characteristics obtained from LDA with complementary information provided by SVD and FrSVD thereby retaining useful information from multiple feature spaces. These fused features from the multiple spaces will help in obtaining the good recognition accuracy.

2.6. Illumination and Pose Variations

Illumination and facial pose variations represent the major challenge in face recognition since they substantially modify the observed appearance of an individual while the underlying identity remains unchanged. Illumination variations alter the intensity distribution over different facial regions, whereas pose changes modify the visible facial structure and the relative arrangement of facial components.

Georghiades et al. [21] Zhang et al. [14] developed illumination-cone models for representing facial appearance under different lighting and pose conditions. This study is mostly focused on significant variations in appearance-based features due to illumination and pose variations in face recognition. Similar type of work was produced by Joseph [16] where face recognition is investigated beyond frontal facial views while considering several sources of variation, including pose, illumination, expression, aging, and occlusion. Collectively, these studies demonstrate that a robust recognition system should preserve identity-specific characteristics while reducing the influence of variations that are unrelated to subject identity.

To evaluate the performance, they have used the Extended Yale Face Database B [21] is particularly suitable for investigating illumination robustness because it contains facial images acquired under substantially different lighting conditions. This database provides a challenging evaluation environment to determine whether a feature representation can maintain class discrimination under illumination conditions change significantly.

2.7. RBF Neural Network for Face Classification

For accurate classification of the fused facial features extracted with help of multiple feature selection and representation techniques is expected In facial recognition systems. According to Er et al. [5] RBF is the best suitable neural classifier for facial feature classification since it has an ability to classify the nonlinear facial feature and also found to be universal approximator. Thakur et al. [6] combined PCA with an RBF neural network for face recognition, demonstrating that dimensionality-reduction techniques can be effectively coupled with RBF-based classification. Similarly, in [13,28] examined a face recognition framework combining PCA, Fisher Linear Discriminant-based processing, and RBF neural networks.

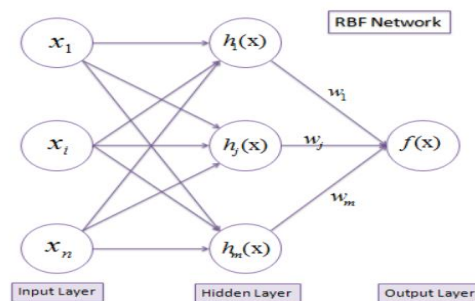


Figure 21.: Architecture of RBF Classifier

A conventional RBF network in figure 2.1. contains consist of input layer, a hidden layer composed of radial basis neurons, and an output layer. The hidden layer transforms the input features into a nonlinear representation according to their distances from selected centers. Such a localized transformation can be advantageous when facial classes exhibit overlapping distributions as a result of illumination, pose, expression, or other appearance variations.

Compared with deeper neural-network architectures, an RBF network can offer a comparatively compact classification structure, especially when the input data have already been reduced through statistical feature extraction. In [23,28] further investigated incremental learning using RBF neural networks for face recognition and demonstrated their potential for adaptive recognition systems. These findings provide support for employing RBF as the nonlinear classification component of the proposed framework.

2.8. Recent Developments in Face Recognition

Recent developments in face recognition have increasingly focused on improving robustness, computational efficiency, and generalization under unconstrained imaging conditions. Deng et al. investigated the use of unrecognizable faces for enhancing face recognition models and examined the relationship between training-data characteristics and recognition performance [24]. Wu et al. [25] studied demographic variations in face

recognition accuracy and emphasized the importance of assessing recognition systems across different population groups. Wang et al. [26] introduced a residual feature decomposition and multi-task learning framework for variation-invariant face recognition. Their results emphasized the importance of constructing representations that preserve identity information while reducing sensitivity to facial variations. Chen et al. [27] investigated confidence-aware RGB-D face recognition using virtual depth synthesis, demonstrating the potential of complementary information and confidence-based integration for enhancing recognition robustness. Khalifa et al. [28] developed an attention-integrated multi-level convolutional neural network for efficient and robust face recognition, reflecting the continuing advancement of feature-learning architectures that seek to improve recognition performance while considering computational requirements.

These studies demonstrate the strong capabilities of modern deep-learning techniques for addressing complex face recognition problems. Nevertheless, deep models often require large-scale training datasets, extensive parameter optimization, and substantial computational resources. In situations where training data are relatively limited or where a simpler and more interpretable processing pipeline is preferred, conventional statistical feature extraction combined with nonlinear classification remains a relevant research direction. This consideration motivates the investigation of hybrid subspace-based approaches capable of integrating complementary feature representations without depending entirely on deep neural architectures.

2.9. Research Gap and Motivation

The existing literature indicates that PCA [1,2] is effective for dimensionality reduction and generate compact facial representations, but it does not explicitly optimize the separation between different identity classes. It also retains the noisy data when projected in image subspace for classification.

LDA [3,11] addresses this limitation by incorporating class-label information and increasing inter-class separability where in SVD [12, 18] provides a compact representation of dominant structural information, whereas fractional SVD-based representation introduces a fractional-order transformation that can provide additional spectral characteristics [15]. Furthermore, RBF neural networks [5, 6, 13, 23,29] have demonstrated their suitability for nonlinear classification of facial feature representations

However, many conventional approaches apply these techniques independently or integrating sequential manner PCA followed by LDA being a representative example. Although PCA-LDA can generate an effective discriminative feature space, such sequential processing may not fully utilize the complementary information available from other structural and spectral representations. Similarly, applying SVD or FrSVD independently may not adequately exploit the class-discriminative characteristics obtained through LDA.

This limitation motivates the development of an optimized framework where complementary feature representations are combined rather than relying exclusively on a single projected feature space. The proposed approach therefore integrates PCA, LDA, SVD, and FrSVD and subsequently combines their complementary representations using a simple weighted feature-fusion strategy. In this framework, PCA is used for initial dimensionality reduction, LDA contributes class-discriminative information, SVD captures dominant structural characteristics, and FrSVD provides an additional fractional-order spectral representation.

The resulting fused feature vector is then provided to an RBF neural network for nonlinear classification. The overall objective is to combine global statistical characteristics, identity-discriminative information, dominant structural properties, and fractional-order spectral information within a unified feature representation while retaining a relatively simple computational architecture.

The proposed framework is evaluated using the Extended Yale Face Database B [21], which contains substantial illumination variations and therefore provides an appropriate benchmark for studying illumination robustness. The experimental investigation examines whether integrating complementary subspace and spectral representations can improve recognition performance compared with conventional PCA-LDA-based approaches. The proposed framework is intended to exploit the complementary strengths of these techniques while avoiding the architectural complexity and extensive data requirements commonly associated with large-scale deep-learning models.

3. Proposed Methodology

The proposed face recognition system employs an integrated framework that combines dimensionality reduction, supervised discriminant analysis, spectral representation, feature fusion, and nonlinear classification using PCA,

LDA, SVD0, FrSVD, and a Radial Basis Function (RBF) neural network respectively. The major objective of proposed methodology is to transform high-dimensional facial images into compact representations, improve the separability of different identity classes, obtain complementary spectral characteristics, and combine the resulting feature representations before classification.

The complete framework as shown in figure 3.1. is organized into four principal stages: (1) image preprocessing and normalization, (2) dimensionality reduction and discriminative feature extraction using PCA and LDA, (3) spectral representation using SVD and FrSVD followed by weighted feature fusion, and (4) nonlinear classification using an RBF neural network.

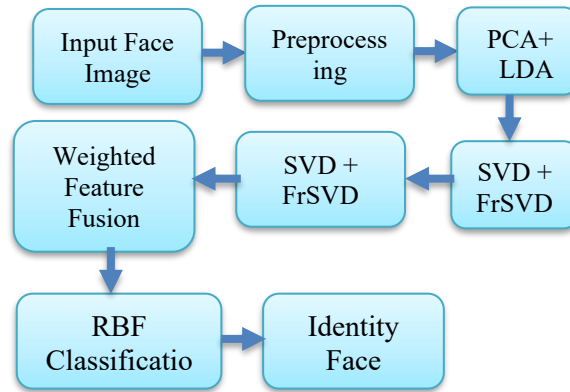


Fig. 3.1. Overall processing sequence of the proposed face recognition framework.

3.1. Preprocessing

The input facial images are initially pre-processed to minimize the influence of illumination variations and establish a consistent representation for the subsequent feature-extraction stages. Histogram equalization is first employed to enhance image contrast and alleviate the effects of non-uniform illumination. The resulting images are converted into grayscale and resized to a predefined spatial resolution. This procedure ensures that every facial sample possesses identical spatial dimensions and can consequently be represented within a common feature space [6].

Let there be N training facial images distributed among C classes. Each grayscale facial image is resized to a fixed dimension of $m \times n$ pixels and subsequently converted into a column vector. Accordingly, the vector representation of the i^{th} image is expressed as

$$x_i \in \mathbb{R}^d, d = m \times n \quad (3.1)$$

Now calculate the mean centred image

$$\mu = \frac{1}{N} \sum_{i=1}^N x_i \quad (3.2)$$

The cantered data matrix is given as to ensure that the PCA operated on zero mean data

$$\tilde{x} = x_i - \mu \quad (3.3)$$

so that the PCA operation is performed using zero-mean data. Mean centering is an important step in PCA because the principal directions are obtained from the covariance structure of the input data. Removing the common mean facial appearance allows the resulting principal components to characterize variations among individual facial samples more effectively.

3.2. Principal Component Analysis (PCA)

A facial image contains a large number of pixel values, making its original representation inherently high-dimensional. Processing these image vectors directly can substantially increase computational requirements and may also increase the possibility of overfitting when the number of available training samples is considerably smaller than the image-space dimensionality (d). To address this issue, Principal Component Analysis (PCA) is initially employed to obtain a lower-dimensional facial representation [1], [2], [9].

Given the mean-centered data matrix X the covariance matrix is calculated as

$$C_T = \frac{1}{N-1} \sum_{i=1}^N \tilde{x}_i \tilde{x}_i^T \quad (3.4)$$

The eigen-decomposition of the covariance matrix is expressed as

$$C u_j = \lambda_j u_j, \quad j = 1, 2, \dots, d \quad (3.5)$$

where λ_j and u_j represent the j^{th} eigenvalue and its corresponding eigenvector, respectively.

The eigenvalues are arranged in descending order as

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d \quad (3.6)$$

The first K eigenvectors associated with the largest eigenvalues are selected to form the PCA projection matrix:

$$U_K = [u_1, u_2, \dots, u_K] \in \mathbb{R}^{d \times K} \quad (3.7)$$

The PCA representation of the i^{th} facial image is obtained by projecting its mean-centered image vector onto the selected principal components:

$$z_i = U_K^T \tilde{x}_i, \quad z_i \in \mathbb{R}^K \quad (3.8)$$

Since

$\tilde{x}_i = x_i - \mu$, the PCA feature vector can equivalently be expressed as

$$z_i = U_K^T (x_i - \mu) \quad (3.9)$$

The PCA transformation retains the principal directions corresponding to the greatest variance in the training data while reducing the dimensionality from d to K . This reduction lowers the computational requirements and produces a compact representation for the subsequent discriminative analysis. Nevertheless, PCA does not use class-label information because it is an unsupervised transformation. Hence, the directions associated with maximum variance are not necessarily those that provide the strongest separation among different face identities. To address this limitation, Linear Discriminant Analysis (LDA) is applied to the PCA-reduced feature representation.

3.3. Linear Discriminant Analysis (LDA)

Although PCA effectively reduces the dimensionality of facial image data, it does not explicitly exploit information concerning the identity classes. Therefore, the feature vectors obtained from PCA are further processed using Linear Discriminant Analysis (LDA) to improve class separability. LDA is a supervised dimensionality-reduction technique that seeks projection directions that increase between-class separation while simultaneously reducing within-class variation [3], [8], [9].

Let the reduced feature vector by PCA corresponding to the i^{th} face image be denoted by

$z_i \in \mathbb{R}^K$, where K represents the number of principal components selected during the PCA stage. Assume that the training samples are divided into C classes, where the c^{th} class contains N_c

The mean vector of the c^{th} class in the PCA-reduced feature space is calculated as

$$\mu_c = \frac{1}{N_c} \sum_{i \in c} z_i \quad (3.10)$$

The global mean of all PCA-reduced training samples is given by

$$\mu_z = \frac{1}{N} \sum_i z_i = 1^N z_i \quad (3.11)$$

The within-class scatter matrix describes the variation of samples with respect to the mean of their corresponding classes. It is defined as

$$S_W = \sum_{c=1}^C \sum_{i \in c} (z_i - \mu_c)(z_i - \mu_c)^T \quad (3.12)$$

The between-class scatter matrix characterizes the variation among the individual class means and is given by

$$S_B = \sum_{c=1}^C N_c (\mu_c - \mu_z)(\mu_c - \mu_z)^T \quad (3.13)$$

The objective of LDA is to determine a projection matrix that maximizes the ratio between the between-class and within-class scatter. This objective can be formulated as

$$J(W) = \frac{|W^T S_B W|}{|W^T S_W W|} \quad (3.14)$$

The optimal discriminant directions are obtained by solving the generalized eigenvalue problem

$$S_B v_j = \lambda_j S_W v_j \quad (3.15)$$

where λ_j denotes the generalized eigenvalue and v_j represents the corresponding generalized eigenvector.

The generalized eigenvalues are arranged in descending order as

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_K \quad (3.16)$$

The first L eigenvectors corresponding to the largest generalized eigenvalues are selected to construct the LDA projection matrix:

$$W_L = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_L] \in \mathbb{R}^{K \times L} \quad (3.17)$$

where L denotes the number of discriminant components retained after LDA.

The LDA representation of the i^{th} facial image is then obtained by projecting its PCA feature vector onto the selected discriminant directions:

$$\mathbf{y}_i = W_L^T \mathbf{z}_i, \quad \mathbf{y}_i \in \mathbb{R}^L \quad (3.18)$$

Thus, the complete PCA-LDA transformation can be represented as

$$\mathbf{y}_i = W_L^T \mathbf{U}_K^T (\mathbf{x}_i - \boldsymbol{\mu}) \quad (3.19)$$

Equation 3.19. indicates that the original facial image vector $\mathbf{x}_i$ is initially centered and projected into the PCA subspace through $\mathbf{U}_K^T$. The resulting PCA feature representation $\mathbf{z}_i$ is subsequently mapped into the LDA discriminant space through W_L^T . Therefore, the dimensionality of the original facial representation is reduced from d to K through PCA and further reduced to L through LDA.

Performing PCA before LDA is particularly useful in face recognition because the original image dimensionality (d) is normally much larger than the number of available training samples. Applying LDA directly to the original image space can therefore result in a singular or poorly conditioned within-class scatter matrix. The preliminary PCA transformation reduces the dimensionality and consequently provides a more stable space in which supervised discriminant analysis can be performed [3], [8], [9].

The resulting feature vector $\mathbf{y}_i$ contains information that is more strongly associated with identity discrimination than the original PCA representation $\mathbf{z}_i$. Specifically, LDA identifies projection directions that increase the separation between different face classes while reducing variations among samples belonging to the same class. However, the LDA representation primarily describes statistical and class-discriminative properties and may not completely preserve additional structural or spectral characteristics present in the original facial image. For this reason, the proposed framework supplements the LDA representation with SVD- and FrSVD-based representations in the following stages.

Accordingly, the LDA output $\mathbf{y}_i$ obtained from the PCA-reduced feature vector $\mathbf{z}_i$ serves as the discriminative component of the final fused representation. The subsequent SVD and FrSVD stages provide complementary spectral information, which is later combined with the LDA representation through the proposed weighted feature-fusion procedure.

3.4. Singular Value Decomposition (SVD)

Although the PCA-LDA representation captures compact and discriminative characteristics of the facial sample, it mainly reflects statistical variance and class-discriminative information. The original facial image also contains structural characteristics that can provide complementary information for recognition. Hence, Singular Value Decomposition (SVD) is applied to the same facial sample to derive a compact spectral representation [12], [18]. The vector representation of the facial image obtained during preprocessing is given by

$$\mathbf{x} * i = [x * i1, x_{i2}, \dots, x_{id}]^T \quad (3.20)$$

For SVD processing, the vector $\mathbf{x}_i$ is reshaped into its corresponding two-dimensional facial image matrix, denoted by $\mathbf{X}_i$. The SVD of the facial image matrix is expressed as

$$\mathbf{X}_i = \mathbf{U}_i \boldsymbol{\Sigma}_i \mathbf{V}_i^T \quad (3.21)$$

where $\mathbf{U}_i$ and $\mathbf{V}_i$ are orthogonal matrices, and $\boldsymbol{\Sigma}_i$ is a diagonal matrix containing the singular values of $\mathbf{X}_i$. The singular values are ordered from the largest to the smallest as

$$\sigma_{i1} \geq \sigma_{i2} \geq \dots \geq \sigma_{ir} \geq 0 \quad (3.22)$$

where

$$r = \min(m, n) \quad (3.23)$$

and σ_{ij} denotes the j^{th} singular value corresponding to the i^{th} facial image.

Accordingly, the singular-value matrix can be written as

$$\boldsymbol{\Sigma} * i = \text{diag}(\sigma * i1, \sigma_{i2}, \dots, \sigma_{ir}) \quad (3.24)$$

The singular values indicate the relative contribution of the associated orthogonal components to the facial image. In general, the larger singular values capture the principal structural content of the image, while the components

associated with smaller singular values make progressively smaller contributions to its overall representation [12], [18].

To construct a compact SVD representation, only the first k dominant singular values are retained, where $k < r$. The SVD feature vector corresponding to the (i) -th facial image is therefore defined as

$$\mathbf{f}_i^{SVD} * i = [\sigma * i1, \sigma_{i2}, \dots, \sigma_{ik}]^T \quad (3.25)$$

The parameter k specifies the number of dominant spectral components preserved in the SVD representation. Selecting only these components reduces the dimensionality of the spectral feature space while retaining the principal structural characteristics of the facial image.

The SVD representation can further be understood in terms of a low-rank approximation of the original facial image. By considering the first (k) singular components, the reconstructed approximation of the (i) -th facial image is expressed as

$$\mathbf{X}_i^{(k)} = \sum_{j=1}^k \sigma_{ij} \mathbf{u} * ij \mathbf{v} * ij^T \quad (3.26)$$

where $\mathbf{u} * ij$ and $\mathbf{v} * ij$ represent the j^{th} left and right singular vectors, respectively. This approximation retains the major structural patterns present in the facial image while omitting components associated with relatively smaller singular values.

Thus, the SVD feature vector $\mathbf{f}_i^{SVD}$ provides a spectral representation of the same facial sample that is processed through the PCA–LDA branch. The representations obtained from the two branches can be summarized as

$$\mathbf{f}_i^{PCA-LDA} \quad (3.27)$$

for the PCA–LDA branch, while the original facial image is simultaneously represented through the SVD branch as

$$\mathbf{f}_i^{SVD} \quad (3.28)$$

where $\mathbf{f}_i^{PCA-LDA}$ denotes the discriminative PCA–LDA features and $\mathbf{f}_i^{SVD}$ represents the dominant spectral characteristics extracted using SVD.

The SVD branch therefore contributes information that is complementary to the PCA–LDA representation. While the singular-value features characterize the structural properties of the facial image from a spectral viewpoint, the LDA component primarily emphasizes characteristics that distinguish between different facial classes. Nevertheless, the dominant singular components can also be influenced by illumination, facial expression, and other appearance variations that alter the global image structure. Consequently, the conventional SVD representation alone may not provide sufficient robustness for face recognition in the presence of substantial appearance variations.

To overcome this limitation, a fractional-order transformation of the singular-value representation is introduced in the subsequent stage. The Fractional Singular Value Decomposition (FrSVD) representation adjusts the contribution of the singular-value components using a fractional-order parameter, thereby providing an additional spectral description of the same facial sample. The resulting SVD and FrSVD representations can then be integrated with the LDA feature vector using the proposed weighted feature-fusion strategy.

3.5. Fractional Singular Value Decomposition (FrSVD)

Although conventional SVD offers a compact description of the dominant structural properties of the i^{th} facial image, the influence of each singular component is governed directly by its corresponding singular-value magnitude. The leading singular components can contain important identity-related characteristics; however, they may also be affected by illumination changes, facial expressions, occlusions, and other variations in facial appearance. To introduce greater control over the contribution of individual spectral components, Fractional Singular Value Decomposition (FrSVD) is employed as a complementary representation [15].

The FrSVD operation is applied to the same facial image matrix $\mathbf{X}_i$ considered in the conventional SVD stage. From (3.29), the SVD of the i^{th} facial image is expressed as

$$\mathbf{X}_i = \mathbf{U}_i \Sigma_i \mathbf{V}_i^T \quad (3.29)$$

where Σ_i contains the singular values arranged in descending order as

$$\sigma * i1 \geq \sigma * i2 \geq \dots \geq \sigma_{ir} \geq 0 \quad (3.30)$$

With

$$r = \min(m, n) \quad (3.31)$$

In conventional SVD, the singular values directly determine the relative contribution of the associated spectral components. In FrSVD, a fractional-order parameter is incorporated to alter this contribution. Let (α) represent the fractional-order parameter. The fractional singular-value matrix is then defined as

$$\Sigma_i^{(\alpha)} = \text{diag}(\sigma * i1^\alpha, \sigma * i2^\alpha, \dots, \sigma_{iR}^\alpha) \quad (3.31)$$

where α determines how the relative influence of the singular-value components is adjusted.

By retaining the same orthogonal matrices U_i and V_i , the fractional representation of the i^{th} facial image is obtained as

$$X_i^{(\alpha)} = U_i \Sigma_i^{(\alpha)} V_i^T \quad (3.32)$$

Thus, the fractional transformation changes the singular-value spectrum while preserving the orthogonal structural basis obtained from the decomposition of the original facial image matrix X_i .

As in conventional SVD, a reduced-dimensional FrSVD representation can be obtained by retaining only the first R fractional singular-value components. Accordingly, the FrSVD feature vector for the i^{th} facial image is defined as

$$\mathbf{f}_i = [\sigma * i1^\alpha, \sigma * i2^\alpha, \dots, \sigma * iR^\alpha]^T, \mathbf{f} * i \in \mathbb{R}^R \quad (3.33)$$

where R specifies the number of spectral components retained and α denotes the selected fractional-order parameter.

For comparison, the conventional SVD feature vector obtained in (3.34) is given by

$$\mathbf{s}_i = [\sigma * i1, \sigma * i2, \dots, \sigma_{iR}]^T \quad (3.34)$$

whereas the corresponding FrSVD feature vector is

$$\mathbf{f}_i = [\sigma * i1^\alpha, \sigma * i2^\alpha, \dots, \sigma_{iR}^\alpha]^T \quad (3.35)$$

Therefore, $\mathbf{s}_i$ and $\mathbf{f}_i$ originate from the same i^{th} facial image and the same singular-value decomposition. Their difference lies in the manner in which the spectral components are weighted. The fractional transformation consequently provides an alternative representation of the same spectral information rather than introducing an independent image source.

The parameter α provides an additional mechanism for adjusting the relative importance of the singular components. When the singular values differ considerably in magnitude, applying the fractional transformation changes the relative dominance of the leading components. As a result, the resulting representation can exhibit characteristics that are different from those of the conventional SVD feature vector. This behavior allows FrSVD to complement the conventional SVD representation rather than merely reproducing it [15].

The overall progression of the feature representations generated at the different stages of the proposed framework can be summarized as

$$\mathbf{x} * i \rightarrow \tilde{\mathbf{x}} * i \rightarrow \mathbf{z} * i \rightarrow \mathbf{y} * i \quad (3.36)$$

for the PCA-LDA branch, and

$$\mathbf{x} * i \rightarrow X * i \rightarrow \Sigma * i \rightarrow \mathbf{s} * i \quad (3.37)$$

for the conventional SVD branch. The fractional spectral branch is represented by

$$\mathbf{x} * i \rightarrow X * i \rightarrow \Sigma * i^{(\alpha)} \rightarrow \mathbf{f} * i \quad (3.38)$$

where $\mathbf{y}_i$ denotes the LDA feature vector, $\mathbf{s}_i$ represents the conventional SVD feature vector, and $\mathbf{f}_i$ denotes the FrSVD feature vector.

The three feature representations describe different but complementary properties of the same facial sample. The LDA feature vector $\mathbf{y}_i$ focuses on class-discriminative characteristics obtained after PCA-based dimensionality reduction. In contrast, the SVD feature vector $\mathbf{s}_i$ captures the dominant spectral properties of the facial image, while the FrSVD feature vector $\mathbf{f}_i$ modifies these spectral characteristics through the fractional-order transformation.

The complementary nature of these representations provides the basis for the subsequent feature-fusion stage. Instead of depending solely on the LDA representation or choosing only one of the two spectral representations, the proposed framework integrates LDA, SVD, and FrSVD features through a weighted fusion approach. Prior to fusion, the individual feature vectors are normalized to reduce the effect of differences in their numerical scales and to prevent any single representation from disproportionately influencing the combined feature vector. The normalized feature vectors are denoted by $\hat{\mathbf{y}} * i$, $\hat{\mathbf{s}} * i$, $\hat{\mathbf{f}}_i$ for the LDA, SVD, and FrSVD representations,

respectively. The weighted feature-fusion process used to integrate these complementary representations is presented in the following section.

3.6. Weighted Feature Fusion

The normalized LDA, SVD, and FrSVD feature vectors describe the same facial sample from different and complementary perspectives. The LDA branch focuses primarily on class-discriminative characteristics, while conventional SVD preserves the dominant structural properties of the facial image and FrSVD modifies the corresponding spectral information through a fractional-order transformation. Integrating these complementary representations can consequently provide a more informative feature description than relying on any single feature representation.

For the i^{th} facial image, let the normalized feature vectors be denoted by

$$\hat{\mathbf{y}} * i, \quad \hat{\mathbf{s}} * i, \quad \hat{\mathbf{f}}_i \quad (3.39)$$

where $\hat{\mathbf{y}} * i$, $\hat{\mathbf{s}} * i$, and $\hat{\mathbf{f}}_i$ correspond to the normalized LDA, SVD, and FrSVD feature vectors, respectively. To integrate these three feature representations, a weighted feature-fusion strategy is adopted. The weighted-sum form of the fused representation can be written as

$$\mathbf{h}_i = w_1 \hat{\mathbf{y}} * i + w_2 \hat{\mathbf{s}} * i + w_3 \hat{\mathbf{f}}_i \quad (3.40)$$

where w_1 , w_2 , and w_3 denote the weighting coefficients assigned to the LDA, SVD, and FrSVD representations, respectively. The weighting coefficients are constrained by

$$w_1 + w_2 + w_3 = 1, \quad w_j \geq 0, \quad j = 1, 2, 3. \quad (3.41)$$

Normalization of the individual feature representations prior to fusion is essential because the LDA, SVD, and FrSVD features can have substantially different numerical ranges and statistical distributions. If the features were combined without normalization, a representation having larger numerical values could disproportionately influence the fused feature vector, even when its actual contribution to recognition is not correspondingly greater. The weighting coefficients offer a direct means of regulating the contribution of the three feature representations. An increase in w_1 gives greater emphasis to the LDA features and consequently strengthens the contribution of class-discriminative information. Likewise, increasing w_2 places greater emphasis on the dominant structural characteristics extracted by conventional SVD, whereas a larger value of (w_3) increases the influence of the fractional spectral characteristics provided by FrSVD.

The proposed fusion mechanism is deliberately kept straightforward. Rather than incorporating an additional learning-based fusion architecture or a computationally intensive optimization procedure, predefined weighting coefficients are employed to integrate the complementary feature representations. This approach maintains a relatively simple overall framework while allowing information from the different feature spaces to be jointly utilized.

The fused representation can alternatively be expressed through weighted concatenation as

$$\mathbf{h}_i = [w_1 \hat{\mathbf{y}} * i \parallel w_2 \hat{\mathbf{s}} * i \parallel w_3 \hat{\mathbf{f}}_i] \quad (3.42)$$

where each feature representation is individually scaled before being incorporated into the combined feature space. Under this formulation, the individual descriptors retain their respective information, while their relative influence is determined by the corresponding weighting coefficient. The resulting concatenated representation is then provided as input to the RBF neural network.

The weighted-sum and weighted-concatenation forms should be regarded as distinct interpretations of the fusion operation and selected consistently with the implementation used in the experimental study. In the proposed framework, the three feature vectors are considered complementary descriptors. Therefore, the weighted-concatenation formulation is adopted to preserve the information available in each feature space without requiring the different feature representations to possess identical dimensionality.

Accordingly, the final fused representation corresponding to the (i^{th}) facial sample is defined as

$$\mathbf{h}_i = [w_1 \hat{\mathbf{y}} * i^T, w_2 \hat{\mathbf{s}} * i^T, w_3 \hat{\mathbf{f}}_i^T]^T. \quad (3.43)$$

If the dimensions of the normalized LDA, SVD, and FrSVD feature vectors are (L), (R), and (R), respectively, then the dimensionality of the resulting fused representation is

$$D_f = L + 2R. \quad (3.44)$$

Hence, the fused feature vector $\mathbf{h}_i$ incorporates the discriminative information generated by the PCA–LDA branch together with complementary spectral characteristics obtained from the conventional SVD and FrSVD branches. The resulting representation is subsequently supplied to the RBF neural network for nonlinear classification.

3.6. RBF Neural Network Classification

The final component of the proposed framework is the classification of the fused facial feature representation using a Radial Basis Function (RBF) neural network. The use of an RBF network is appropriate because the preceding PCA, LDA, SVD, FrSVD, and feature-fusion operations transform the original high-dimensional facial image into a compact feature space, where the class distributions may still exhibit nonlinear separation [5], [6], [13], [23].

Let the fused feature vector corresponding to the i^{th} facial image be represented by $\mathbf{h}_i$.

The RBF neural network is composed of three principal layers: an input layer, a hidden layer consisting of radial basis neurons, and an output layer associated with the different face classes. The input layer accepts the fused feature vector $\mathbf{h}_i$, whereas the hidden layer applies a nonlinear transformation according to the distance between the input vector and the center of each radial basis neuron.

For the j^{th} hidden neuron, the radial basis function is defined as

$$\phi_j(\mathbf{h}_i) = \exp\left(-\frac{|\mathbf{h}_i - \mathbf{c}_j|^2}{2\sigma_j^2}\right) \quad (3.45)$$

where $\mathbf{c}_j$ denotes the center of the j^{th} radial basis neuron, σ_j represents its spread parameter, and $|\mathbf{h}_i - \mathbf{c}_j|$ denotes the Euclidean distance between the input feature vector and the corresponding neuron center.

The Gaussian radial basis function provides a strong response when the input feature vector is located close to the associated center. As the distance between the input and the center increases, the activation decreases progressively. In this manner, the hidden layer transforms the fused feature representation into a nonlinear feature space, potentially making the feature distributions of different facial identities more distinguishable.

If the RBF network contains M hidden neurons, the hidden-layer output corresponding to the i^{th} facial image can be represented as

$$\boldsymbol{\phi}_i = [\phi_1(\mathbf{h}_i), \phi_2(\mathbf{h}_i), \dots, \phi_M(\mathbf{h}_i)]^T. \quad (3.46)$$

The response of the $(c^{\wedge}\{th\})$ output class node is then given by

$$o_{ic} = \sum_{j=1}^M w_{jc} \phi_j(\mathbf{h}_i) + b_c \quad (3.47)$$

where w_{jc} denotes the connection weight between the j^{th} hidden neuron and the c^{th} output node, while b_c represents the bias associated with that output node.

The identity predicted for the input facial image is determined by selecting the output node that produces the largest response:

$$\hat{c}_i = \arg \max_c o_{ic}. \quad (3.48)$$

Thus, the RBF network performs nonlinear classification within the fused feature space. This characteristic is particularly beneficial when samples belonging to different facial identities have partially overlapping feature distributions as a consequence of illumination changes, pose variations, facial expressions, occlusions, or other appearance-related factors.

3.7. Overall Algorithm of the Proposed Framework

The complete sequence of operations in the proposed face-recognition framework is summarized in Algorithm 1.

Algorithm 1: PCA–LDA–SVD–FrSVD with Weighted Feature Fusion and RBF Classification

Input: Training images $\mathbf{x} * i$, test image $\mathbf{x} * test$, PCA dimension K , LDA dimension L , SVD dimension R , fractional order (α), fusion weights (w_1, w_2, w_3) , and RBF parameters.

Output: Predicted identity $\hat{c}_{test}$

1. **Preprocessing:** Convert all training and test facial images to grayscale, perform histogram equalization, resize the images to the required dimensions, and vectorize the resulting images.
2. **PCA projection:** Calculate the training-data mean vector, covariance matrix, and PCA projection matrix. Project the centered training and test samples onto the first K principal components to obtain the reduced PCA representations.

3. **LDA projection:** Using the PCA-derived training features together with their corresponding class labels, calculate the within-class scatter matrix S_W and between-class scatter matrix S_B . Solve the generalized eigenvalue problem

$$S_B \mathbf{v} = \lambda S_W \mathbf{v}$$

and retain the first L discriminant directions. The resulting LDA representations are denoted by $\mathbf{y} * i$ for the training samples and $\mathbf{y} * test$ for the test sample.

4. **SVD and FrSVD extraction:** Reshape each pre-processed facial vector into its corresponding two-dimensional image matrix and perform the decomposition

$$X_i = U_i \Sigma_i V_i^T.$$

Retain the first (R) singular values to construct the conventional SVD feature vector

$$\mathbf{s}_i = [\sigma_{i1}, \sigma_{i2}, \dots, \sigma_{iR}]^T.$$

Using the selected fractional order (α), construct the corresponding FrSVD feature vector as

$$\mathbf{f}_i = [\sigma_{i1}^\alpha, \sigma_{i2}^\alpha, \dots, \sigma_{iR}^\alpha]^T.$$

The same procedure is applied to the test image to obtain $\mathbf{s} * test$ and $\mathbf{f} * test$.

5. **Feature normalization and fusion:** Normalize the LDA, SVD, and FrSVD feature vectors to reduce differences in their numerical scales. The normalized representations are then combined using the predefined fusion weights:

$$\mathbf{h}_i = [w_1 \hat{\mathbf{y}}_i^T, w_2 \hat{\mathbf{s}}_i^T, w_3 \hat{\mathbf{f}}_i^T]^T$$

The same normalization parameters and fusion weights are used for the test representation.

6. **RBF classification:** Train the RBF neural network using the fused training representations $\mathbf{h}_i$ and their corresponding identity labels. The hidden neurons apply the Gaussian radial basis function

$$\phi_j(\mathbf{h}_i) = \exp\left(-\frac{\|\mathbf{h}_i - \mathbf{c}_j\|^2}{2\sigma_j^2}\right)$$

The resulting hidden-layer activations are used to determine the class-specific output responses.

7. **Test-image classification:** Apply the learned PCA, LDA, SVD, and FrSVD transformations, together with the previously determined normalization parameters and fusion weights, to the test image to obtain

$$\mathbf{h}_{test} = [w_1 \hat{\mathbf{y}} * test^T, w_2 \hat{\mathbf{s}} * test^T, w_3 \hat{\mathbf{f}}_{test}^T]^T.$$

The fused test representation is then passed through the trained RBF network. The predicted identity is assigned according to

$$\hat{C}_{test} = \arg \max_c o_{test,c},$$

where $o_{test,c}$ denotes the RBF output corresponding to class c .

The above sequence completes the proposed face-recognition process, beginning with preprocessing and PCA-based dimensionality reduction, followed by LDA-based discrimination, SVD and FrSVD spectral feature extraction, weighted feature fusion, and finally nonlinear classification using the RBF neural network.

4. Experimental Results and Comparative Analysis

This section presents the experimental evaluation of the proposed PCA–LDA–SVD–FrSVD framework with weighted feature fusion and RBF neural network classification. The experiments are designed to evaluate the effectiveness of the proposed method for face recognition under challenging illumination and facial appearance variations. The performance of the proposed framework is analyzed using the Extended Yale Face Database B and is compared with conventional feature-extraction and classification approaches.

The experimental evaluation focuses on four main aspects: (i) the recognition performance of the proposed framework, (ii) the contribution of the individual feature representations, (iii) the effect of weighted feature fusion, and (iv) comparison with conventional and related face recognition methods.

A. Data Set Used for Experimentation

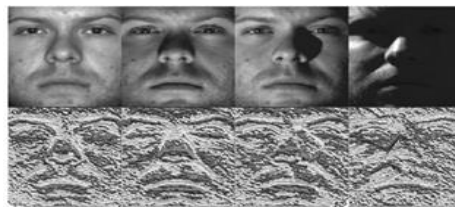
The Extended Yale Face Database B is employed to assess the performance of the proposed face recognition framework [21]. This database is a commonly used benchmark for studying illumination-robust face recognition, as it provides facial images of several subjects captured under a wide range of lighting conditions.

The dataset consists of facial samples from different individuals recorded under various illumination settings. These changes in illumination can substantially alter the visual appearance of a face while the underlying subject identity remains unchanged. Consequently, the database provides an appropriate experimental setting for examining whether the proposed combination of discriminative and spectral feature representations can preserve class separability in the presence of illumination variations.

Prior to feature extraction, every image undergoes the preprocessing procedure presented in Section 3. Histogram equalization is first performed to enhance image contrast, after which the images are converted to grayscale and resized to a uniform spatial resolution.



(a) Yale Database



(b) Fishers Faces



(c) Eigen Faces

B. Performance Evaluation

Table 1: Performance Comparison

Method	Accuracy (%)
PCA	85.2
PCA + LDA	91.4
PCA + LDA + SVD	93.1
PCA + LDA + FrSVD + RBF	94.6
PCA + LDA + FrSVD + RBF with weighted fusion	96.5

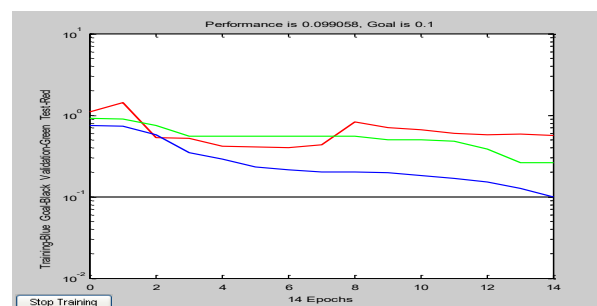


Figure-4.1-Comparative Analysis using FrSVD and RBF

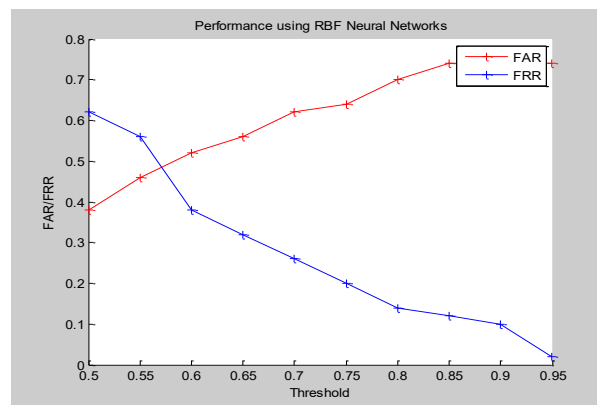


Figure-4.2- Comparative Analysis of FRR and FAR for FrSVD and RBF with Weighted fusion

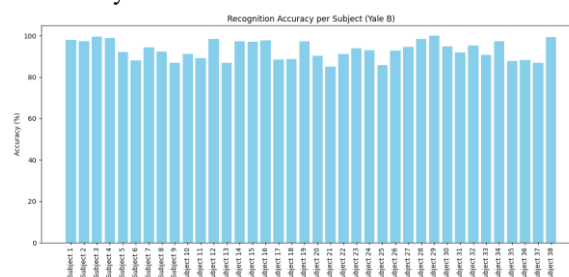


Figure 4.3: Recognition Accuracy of Yale Database

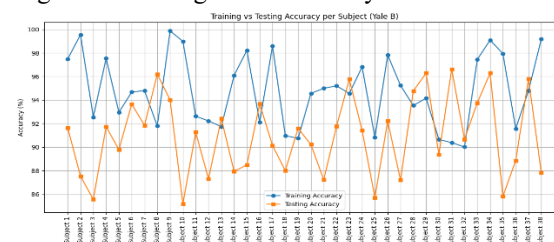


Figure 4.4: Recognition Accuracy of Yale Database

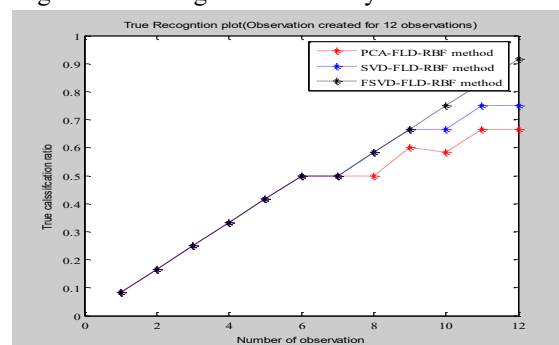


Figure-4.4 Comparison of three Different Methods

5 . Conclusion

This study proposed a hybrid face recognition framework integrating PCA, LDA, SVD, FrSVD, weighted feature fusion, and an RBF neural network. PCA reduces dimensionality, LDA enhances class discrimination, while SVD and FrSVD provide complementary structural and fractional spectral information. The fused representation is classified using an RBF network to obtain nonlinear decision boundaries. Experimental evaluation on the Extended Yale Face Database B achieved a recognition accuracy of 96.5%, demonstrating the effectiveness of the proposed approach under illumination variations.

However, the study is limited by evaluation on a single database and the use of predefined fusion weights and fractional-order parameters. The method may require further validation under severe pose, occlusion, aging, and unconstrained environmental conditions.

Future work will focus on evaluating multiple challenging datasets, optimizing fusion weights and FrSVD parameters automatically, incorporating local or deep features, and improving computational efficiency. Adaptive learning, real-time implementation, and robust evaluation across diverse demographic and environmental conditions are also promising directions.

References

- [1] M. Turk and A. Pentland, "Eigenfaces for recognition," *J. Cogn. Neurosci.*, vol. 3, no. 1, pp. 71–86, 1991, doi: 10.1162/jocn.1991.3.1.71.
- [2] Y. Zhang and L. Yan, "A fast face recognition based on image gradient compensation for feature description," *Multimedia Tools and Applications*, vol. 81, pp. 26015–26034, 2022, doi: 10.1007/s11042-022-12804-4. [Online]. Available: <https://link.springer.com/article/10.1007/s11042-022-12804-4>.
- [3] P. N. Belhumeur, J. P. Hespanha, and D. J. Kriegman, "Eigenfaces vs. Fisherfaces: Recognition using class specific linear projection," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 19, no. 7, pp. 711–720, Jul. 1997.
- [4] R. Chellappa, C. L. Wilson, and S. Sirohey, "Human and machine recognition of faces: A survey," *Proc. IEEE*, vol. 83, no. 5, pp. 705–741, May 1995.
- [5] M. J. Er, S. Wu, J. Lu, and H. L. Toh, "Face recognition with radial basis function (RBF) neural networks," *IEEE Trans. Neural Netw.*, vol. 13, no. 3, pp. 697–710, May 2002.
- [6] S. Thakur, J. K. Sing, D. K. Basu, M. Nasipuri, and M. Kundu, "Face recognition using principal component analysis and RBF neural networks," in *Proc. 1st Int. Conf. Emerging Trends in Engineering and Technology (ICETET)*, 2008, pp. 695–700, doi: 10.1109/ICETET.2008.104.
- [7] "TIFd-FR: Trends, issues and future directions of feature extraction in face recognition," *Procedia Computer Science*, vol. 235, pp. 1386–1398, 2024, doi: 10.1016/j.procs.2024.04.130. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1877050924008068>
- [8] Y. Bakhshi, S. Kaur, and P. Verma, "A study based on various face recognition algorithms," *Int. J. Comput. Appl.*, vol. 129, no. 13, pp. 16–20, 2015.
- [9] K. I. Kim, K. Jung, and H. J. Kim, "Face recognition using kernel principal component analysis," *IEEE Signal Process. Lett.*, vol. 9, no. 2, pp. 40–42, Feb. 2002, doi: 10.1109/97.991133.
- [10] S. F. Hafez, M. M. Selim, and H. H. Zayed, "2D face recognition system based on selected Gabor filters and linear discriminant analysis LDA," *arXiv preprint arXiv:1503.03741*, 2015.
- [11] V. A. Kannamed, B. S. Murthy, and N. Subramanyam, "Performance study of LDA and KFA for Gabor based face recognition system," *Procedia Computer Science*, vol. 57, pp. 960–969, 2015, doi: 10.1016/j.procs.2015.07.493.
- [12] M. Deswal, N. Kumar, and N. Rathi, "Face recognition based on singular value decomposition linear discriminant analysis method," *Int. J. Eng. Innov. Technol.*, vol. 3, no. 12, pp. 192–195, 2014.
- [13] K. C. Jondhale and L. M. Waghmare, "Improvement in PCA performance using FLD and RBF neural networks for face recognition," in *Proc. 3rd Int. Conf. Emerging Trends in Engineering and Technology (ICETET)*, 2010, pp. 500–505.
- [14] W. Zhang, X. Zhao, J.-M. Morvan, and L. Chen, "Improving shadow suppression for illumination robust face recognition," in *Proc. IEEE Int. Conf. Computer Vision (ICCV)*, 2017.
- [15] J. Liu, S. Chen, and X. Tan, "Fractional order singular value decomposition representation for face recognition," *Pattern Recognition*, vol. 41, no. 1, pp. 378–395, 2008, doi: 10.1016/j.patcog.2007.03.027.
- [16] M. Joseph and K. Elleithy, "Beyond frontal face recognition," *IEEE Access*, vol. 11, pp. 26850–26861, 2023, doi: 10.1109/ACCESS.2023.3258444.
- [17] O. I. Y. D. Bashi, "Face recognition based on PCA, LBP and SVM techniques," *Eng. Technol. J.*, vol. 33, 2015.
- [18] J. Lu and Y. Zhao, "Dominant singular value decomposition representation for face recognition," *Signal Processing*, vol. 90, no. 6, pp. 2087–2093, Jun. 2010, doi: 10.1016/j.sigpro.2009.11.028.
- [19] E. J. Kansa, "Radial basis functions: Achievements and challenges," in *Boundary Elements and Other Mesh Reduction Methods XXXVII*, vol. 61, pp. 3–22, 2015.
- [20] AT&T Laboratories Cambridge, "The Database of Faces," Cambridge, U.K. [Online]. Available: <https://www.cl.cam.ac.uk/research/dtg/attarchive/facedatabase.html>

- [21] A. S. Georghiades, P. N. Belhumeur, and D. J. Kriegman, "From few to many: Illumination cone models for face recognition under variable lighting and pose," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 23, no. 6, pp. 643–660, Jun. 2001.
- [22] J. Yang, J. Yang, D. Zhang, and J.-F. Lu, "Feature fusion: Parallel strategy vs. serial strategy," *Pattern Recognition*, vol. 36, no. 6, pp. 1369–1381, Jun. 2003, doi: 10.1016/S0031-3203(02)00262-5.
- [23] Y. W. Wong, K. P. Seng, and L.-M. Ang, "Radial basis function neural network with incremental learning for face recognition," *IEEE Trans. Syst., Man, Cybern. B, Cybern.*, vol. 41, no. 4, pp. 940–949, Aug. 2011, doi: 10.1109/TSMCB.2010.2101591.
- [24] S. Deng, Y. Xiong, M. Wang, W. Xia, and S. Soatto, "Harnessing unrecognizable faces for improving face recognition," in *Proc. IEEE/CVF Winter Conf. Applications of Computer Vision (WACV)*, 2023, pp. 3424–3433.
- [25] H. Wu, V. Albiero, K. S. Krishnapriya, M. C. King, and K. W. Bowyer, "Face recognition accuracy across demographics: Shining a light into the problem," in *Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2023, pp. 1041–1050.
- [26] A. Haider, G. Wu, I. Spence, and H. Wang, "Residual feature decomposition and multi-task learning-based variation-invariant face recognition," *Neural Computing and Applications*, vol. 36, pp. 20147–20166, 2024, doi: 10.1007/s00521-024-10234-x.
- [27] Z. Chen, M. Wang, W. Deng, H. Shi, D. Wen, Y. Zhang, X. Cui, and J. Zhao, "Confidence-aware RGB-D face recognition via virtual depth synthesis," in *Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2024, pp. 1481–1489.
- [28] A. Khalifa, A. A. Abdelrahman, T. Hempel, and A. Al-Hamadi, "Towards efficient and robust face recognition through attention-integrated multi-level CNN," *Multimedia Tools and Applications*, vol. 84, pp. 12715–12737, 2025, doi: 10.1007/s11042-024-19521-0.