

基于欧拉表示的线性判别分析

赵永胜

重庆师范大学计算机与信息科学学院, 重庆

收稿日期: 2025年1月15日; 录用日期: 2025年2月17日; 发布日期: 2025年2月26日

摘要

线性判别分析(LDA)是一种经典的线性降维方法, 广泛应用于统计学和机器学习领域, 用于提取数据中的判别特征。然而, 随着数据规模的迅速扩大, 大量数据呈现出复杂的非线性特征, 使得传统依赖欧氏距离的LDA在处理此类数据时难以充分捕捉复杂的数据结构。为了解决这一问题, 本文提出了一种基于欧拉表示的线性判别分析模型(Euler-LDA)。Euler-LDA结合欧拉表示的优势, 将数据映射到复数空间。通过利用复数的几何特性, 更精确地捕捉数据的非线性关系, 从而显著提高算法在复杂数据分布场景下的适用性和分类性能。此外, 这种方法在处理复杂数据分布时, 展现出较强的适用性与优越的分类性能。此外, 通过进一步拉大类间距离与缩小类内距离, Euler-LDA有效提升了特征提取的准确性。通过这些机制, 欧拉表示能够更加高效地应对非线性分布的数据, 提供比传统欧氏距离更可靠和精确的相似性度量。在多个数据库上的对比实验结果表明, 该算法识别率显著优于传统LDA及其改进方法。

关键词

线性判别分析, 欧拉表示法, 数据降维, 机器学习

Linear Discriminant Analysis Based on Euler Representation

Yongsheng Zhao

College of Computer and Information Science, Chongqing Normal University, Chongqing

Received: Jan. 15th, 2025; accepted: Feb. 17th, 2025; published: Feb. 26th, 2025

Abstract

Linear Discriminant Analysis (LDA) is a classic linear dimensionality reduction method widely used in statistics and machine learning for extracting discriminative features from data. However, as data scales rapidly expand, a large amount of data exhibits complex nonlinear characteristics, making it difficult for traditional LDA, which relies on Euclidean distance, to fully capture intricate data

文章引用: 赵永胜. 基于欧拉表示的线性判别分析[J]. 计算机科学与应用, 2025, 15(2): 142-152.

DOI: 10.12677/csa.2025.152041

structures when processing such data. To address this issue, this paper proposes a Linear Discriminant Analysis model based on Euler representation (Euler-LDA). Euler-LDA leverages the advantages of Euler representation to map data into a complex space. By utilizing the geometric properties of complex numbers, it more accurately captures the nonlinear relationships in the data, thereby significantly improving the algorithm's applicability and classification performance in scenarios with complex data distributions. Furthermore, this method demonstrates strong applicability and superior classification performance when dealing with complex data distributions. Additionally, by further increasing the inter-class distance and reducing the intra-class distance, Euler-LDA effectively enhances the accuracy of feature extraction. Through these mechanisms, Euler representation can more efficiently handle nonlinearly distributed data, providing a more reliable and precise similarity measure than traditional Euclidean distance. Comparative experimental results on multiple databases show that the recognition rate of this algorithm is significantly superior to that of traditional LDA and its improved methods.

Keywords

Linear Discriminant Analysis, Euler Representation, Data Dimensionality Reduction, Machine Learning

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

在信息高速发展的时代,海量图像数据为推动图像识别技术的进步提供了丰富的资源。然而,庞大的数据量也带来了维度灾难问题,有效处理多维数据成为降维(Dimensionality Reduction, DR) [1]领域的核心挑战。在众多降维算法中,主成分分析(Principal Component Analysis, PCA) [2]和局部保持投影(Locality Preserving Projections, LPP) [3]备受关注。PCA旨在寻找一种线性投影,最大化样本协方差或最小化重建误差;而LPP则通过构建Laplacian图矩阵计算投影矩阵,将原始数据点映射到新的子空间,同时保持原始数据与映射数据局部结构的一致性,从而优化邻域信息的保留。然而,这两种方法未能充分提取数据中的判别信息。

线性判别分析(Linear Discriminant Analysis, LDA) [4]是一种常用的监督型降维技术,旨在提取数据的判别特征,通过寻求一种线性映射,在降低维度的同时最大化类间散度与类内散度的比值。尽管PCA、LDA和LPP均属于线性降维方法,但它们假设数据是线性可分的,而实际数据通常表现为高维、复杂的非线性分布,不符合这一假设。为了解决这一问题,研究人员提出了多种基于核的降维方法,如核主成分分析(Kernel Principal Component Analysis, KPCA) [5]和核线性判别分析(Kernel Discriminant Analysis, KDA) [6]。这些方法通过将数据映射到高维特征空间,以线性方式捕捉非线性关系,从而在低维空间中更好地保留数据的非线性特征。

文献[7] [8]表明,一种基于余弦距离的新型度量方法能够很好地接近理想的鲁棒核,这种方法被命名为欧拉表示(Euler Representation)。欧拉表示不仅能够实现对理想鲁棒核的近似,还能够显著拉大类间和类内距离的差异。因此,诸如Euler SRC [9]、Euler PCA [10]和Euler LPP [11]等算法通过采用欧拉表示来进行距离或相似性计算,大幅提升了处理非线性复杂数据的能力。针对传统方法在处理非线性分布数据时的不足,Euler-LDA结合欧拉表示的优势,将数据映射到复数空间。通过利用复数的几何特性,更精确

地捕捉数据的非线性关系，从而显著提高算法在复杂数据分布场景下的适用性和分类性能。此外，该方法设计了一种改进的目标损失函数，通过最大化类间散布矩阵与类内散布矩阵的差异来优化分类性能。这一优化策略不仅提高了对数据判别信息的提取能力，还避免了矩阵奇异性问题，从而有效缓解了小样本问题带来的影响。通过这些改进，Euler-LDA 显著提升了在非线性和高维低样本场景中的鲁棒性与判别能力。

本文的其余内容安排如下：第 2 节回顾了经典线性判别分析(LDA)及欧拉表示的相关基础知识；第 3 节详细阐述了提出的新模型；第 4 节聚焦于目标函数的设计；第 5 节通过公共数据集上的实验验证了所提算法的性能；最后，第 6 节对全文进行了总结。

2. 理论知识

本文提出了 Euler-LDA (基于欧拉表示的线性判别分析)。因此，文章首先回顾了 LDA 的基本理论，然后介绍了欧拉表示法。这为我们提出的新方法奠定了坚实的理论基础。

2.1. 线性判别分析

线性判别分析(LDA)是统计学和机器学习领域的一种经典线性降维技术。用于从数据中提取判别信息。基于类别间和类内的分散性质，LDA 寻找从高维空间到更低维空间的线性映射，并最大化类别之间的可分离性。

假设在 $\mathbb{R}^D$ 中有 l 个原始样本 $\mathbf{X} = [x_1, x_2, \dots, x_l]$ ，LDA 目标寻找一个最优投影矩阵 $\mathbf{W}$ ，将高维空间的数据投影到较低维度的空间时，最大化不同类别之间的距离，同时最小化同一类别之间的距离，该方法旨在寻求线性映射：

$$\mathbf{y}_i = \mathbf{W}^T \mathbf{x}_i. \quad (1)$$

对于 LDA 方法，首先计算每个类别的样本中心 $\mathbf{M}_o$ ， $\mathbf{M}_o$ 为第 o 类的样本的均值向量，用来测量类内样本和类间样本的距离，其中 k 为第 o 类样本的数量。

$$\mathbf{M}_o = \frac{1}{k} \sum_{i=1}^k \mathbf{x}_i. \quad (2)$$

将类别标签表示为 $c_i = [1, 2, \dots, k_c]$ ， k_c 是样本的类别数， $\mathbf{M}$ 是所有样本的平均向量。样本的类内散布矩阵 $\mathbf{S}_w$ 和类间散布矩阵 $\mathbf{S}_b$ 分别表示类间散度矩阵和类内散度矩阵，可计算为：

$$\begin{aligned} \mathbf{S}_w &= \sum_{o=1}^{k_c} \sum_{i \in c_i} (\mathbf{x}_i - \mathbf{M})(\mathbf{x}_i - \mathbf{M}_o)^T \\ \mathbf{S}_b &= \sum_{o=1}^{k_c} (\mathbf{M}_o - \mathbf{M})(\mathbf{M}_o - \mathbf{M})^T \end{aligned} \quad (3)$$

考虑线性映射 $\mathbf{W} \in \mathbb{R}^{D \times d}$ ，其中 D 和 d 分别是变换之前和之后数据的维度，LDA 由 Fisher 判别准则针对多分类问题提出的，旨在最大化损失函数：

$$J(\mathbf{W}) = \arg \max_{\mathbf{W}} \frac{\text{tr}(\mathbf{W}^T \mathbf{S}_b \mathbf{W})}{\text{tr}(\mathbf{W}^T \mathbf{S}_w \mathbf{W})} \quad (4)$$

2.2. 欧拉表示法

欧氏距离假设数据的分类边界是线性的，而实际应用中许多数据集的分类边界是非线性的，欧氏距离无法有效处理这些非线性分布的数据，在面对非线性分布时可能会失效，导致距离计算结果不准确。

核技术可以捕获特征的非线性相似性以抑制异常值，我们提出了一种基于余弦距离的新距离度量，该距离度量可以近似理想的鲁棒内核。欧拉表示是一种用于改进距离度量的数学方法，通过将余弦距离形成一种灵活且高效的非线性相似性度量方式。在欧拉表示中，两个向量之间的相似性通过余弦距离来描述。给定两个任意向量 $\mathbf{x}_i$ 和 $\mathbf{x}_b \in \mathbb{R}^D$ ，它们之间的余弦距离定义为：

$$d(\mathbf{x}_i, \mathbf{x}_b) = \sum_{a=1}^D \left\{ 1 - \cos(\pi(\mathbf{x}_i(a) - \mathbf{x}_b(a))) \right\} \quad (5)$$

其中 $\mathbf{x}_i(a)$ 和 $\mathbf{x}_b(a)$ 分别是 $\mathbf{x}_i$ 和 $\mathbf{x}_b$ 的第 a 个分量，方程(5)本质上是一个傅里叶余弦级数的表示。其核心思想是利用傅里叶余弦级数的逼近特性，为复杂数据分布提供更准确的相似性描述，从而满足非线性关系度量的需求。由于傅里叶定理指出任何连续函数都可以用一系列正弦曲线来描述，因此核函数可以在傅里叶空间中展开并进行近似。通过截取有限数量的傅里叶余弦分量，将复杂的核函数映射为若干函数形式的线性组合。这种方法不仅有效简化了核函数的复杂性，还提升了对复杂数据分布的建模效率和相似性度量能力。

通过简单代数，(5)变为：

$$d(\mathbf{x}_i, \mathbf{x}_b) = \sum_{a=1}^D \left\{ 1 - \cos(\pi(\mathbf{x}_i(a) - \mathbf{x}_b(a))) \right\} = \left\| \frac{1}{\sqrt{2}} (e^{i\pi\mathbf{x}_i} - e^{i\pi\mathbf{x}_b}) \right\|^2 = \|\mathbf{g}_i - \mathbf{g}_b\|^2 \quad (6)$$

其中，

$$\mathbf{g}_i = \frac{1}{\sqrt{2}} \begin{bmatrix} e^{i\pi\mathbf{x}_i(1)} \\ \cdot \\ \cdot \\ e^{i\pi\mathbf{x}_i(D)} \end{bmatrix} = \frac{1}{\sqrt{2}} e^{i\pi\mathbf{x}_i} \quad (7)$$

$\mathbf{g}_i$ 称为 $\mathbf{x}_i$ 的欧拉表示。这样， $\mathbf{x}_i$ 和 $\mathbf{x}_b$ 之间的余弦距离就可以看作 $\mathbf{g}_i$ 和 $\mathbf{g}_b$ 之间的欧氏距离，欧拉表示能够捕获传统线性距离无法体现的复杂相似性关系。

3. 基于欧拉表示的线性判别分析的算法设计

假设在 $\mathbb{R}^D$ 空间中存在原始样本 $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_l]$ ，令 $c_o = [1, 2, \dots, k_c]$ 表示类别标签，其中 k_c 是类别的总数。每个类别的样本中心 $\mathbf{M}_o$ ， $\mathbf{M}_o$ 为第 o 类的样本的均值向量，用来衡量类内样本的相对距离，其中 k 为第 o 类样本的数量。

$\mathbf{M}$ 表示所有样本的平均向量，用来衡量类间样本的相对距离。由公式(3)，样本的类内散布矩阵 $\mathbf{S}_w$ 和类间散布矩阵 $\mathbf{S}_b$ 分别表示类间距离和类内距离，可计算为：

$$\begin{aligned} \mathbf{S}_w &= \sum_{o=1}^{k_c} \sum_{i \in C_i} (\mathbf{x}_i - \mathbf{M}_o)(\mathbf{x}_i - \mathbf{M}_o)^T \\ \mathbf{S}_b &= \sum_{o=1}^{k_c} (\mathbf{M}_o - \mathbf{M})(\mathbf{M}_o - \mathbf{M})^T \end{aligned} \quad (8)$$

将 $\mathbf{x}_i$ 投影到一个复数空间中得到 $\underline{\mathbf{x}}_i$ 。

$$\underline{\mathbf{x}}_i = \frac{1}{\sqrt{2}} \begin{bmatrix} e^{i\pi\mathbf{x}_i(1)} \\ \cdot \\ \cdot \\ e^{i\pi\mathbf{x}_i(D)} \end{bmatrix} = \frac{1}{\sqrt{2}} e^{i\pi\mathbf{x}_i} \quad (9)$$

然后，在复空间中基于欧拉表示计算每个类别在复数空间中的样本中心 $\underline{\mathbf{M}}_o$ ， $\underline{\mathbf{M}}_o$ 为第 o 类的样本的均值向量，用来衡量在复数空间中的类内样本和类间样本的相对距离，其中 k 为第 o 类样本的数量。 $\underline{\mathbf{M}}$ 表示所有样本的平均向量。

$$\begin{aligned}\underline{\mathbf{M}}_o &= \frac{1}{k} \sum_{i=1}^k \mathbf{x}_i \\ &= \frac{1}{k} \sum_{i=1}^k \frac{1}{\sqrt{2}} e^{i\pi \mathbf{x}_i} \\ &= \frac{1}{\sqrt{2}} e^{i\pi \underline{\mathbf{M}}_o}\end{aligned}\quad (10)$$

在新的欧拉表示中，样本 $\mathbf{x}_i$ 和样本中心 $\mathbf{M}_o$ 之间的余弦距离可以被定义为：

$$\begin{aligned}d(\mathbf{x}_i, \mathbf{M}_o) &= \sum_{o=1}^{k_c} \sum_{i \in C_o} \left\{ 1 - \cos(\pi(\mathbf{x}_i - \mathbf{M}_o)) \right\} \\ &= \left\| \frac{1}{\sqrt{2}} (e^{i\pi \mathbf{x}_i} - e^{i\pi \mathbf{M}_o}) \right\|^2 \\ &= \|\underline{\mathbf{x}}_i - \underline{\mathbf{M}}_o\|^2\end{aligned}\quad (11)$$

$\underline{\mathbf{x}}_i$ 称为 $\mathbf{x}_i$ 的欧拉表示， $\underline{\mathbf{M}}_o$ 称为 $\mathbf{M}_o$ 的欧拉表示。这样， $\mathbf{x}_i$ 和 $\mathbf{M}_o$ 之间的余弦距离就可以看作 $\underline{\mathbf{x}}_i$ 和 $\underline{\mathbf{M}}_o$ 之间的欧氏距离。使用新的欧拉表示，计算新的类内散度矩阵 $\mathbf{S}_w$ 得：

$$\mathbf{S}_w = \sum_{o=1}^{k_c} \sum_{i \in C_o} (\underline{\mathbf{x}}_i - \underline{\mathbf{M}}_o)(\underline{\mathbf{x}}_i - \underline{\mathbf{M}}_o)^T \quad (12)$$

同理， $\underline{\mathbf{M}}$ 称为 $\mathbf{M}$ 的欧拉表示。 $\underline{\mathbf{M}}_o$ 和 $\mathbf{M}$ 之间的余弦距离就可以看作 $\underline{\mathbf{M}}_o$ 和 $\underline{\mathbf{M}}$ 之间的欧氏距离。使用新的欧拉表示，计算新的类内散度矩阵 $\mathbf{S}_b$ 得：

$$\mathbf{S}_b = \sum_{o=1}^{k_c} (\underline{\mathbf{M}}_o - \underline{\mathbf{M}})(\underline{\mathbf{M}}_o - \underline{\mathbf{M}})^T \quad (13)$$

4. 目标函数设计

LDA 的目标通常是最大化类间散布矩阵 $\mathbf{S}_b$ 和类内散布矩阵 $\mathbf{S}_w$ 的比值来实现最优分类。然而，在小样本问题中，样本量不足可能导致类内散布矩阵 $\mathbf{S}_w$ 奇异，从而无法进行矩阵求逆，影响 LDA 的有效性。在求解的过程中需要计算 $\mathbf{S}_w$ 的逆矩阵，会遇到小样本问题。这一问题在求解过程中尤为突出，因为计算类内散布矩阵的逆矩阵是关键步骤之一。为解决这一问题，我们采用了一种基于 LDA 的改进形式，其目标函数旨在最大化以下损失函数：

$$J = \arg \max_{\mathbf{W}} \text{tr}(\mathbf{W}^T (\mathbf{S}_b - \mathbf{S}_w) \mathbf{W}) \quad (14)$$

新的损失函数就避免了类内散布矩阵 $\mathbf{S}_w$ 的求逆，我们能够更好地应对样本数量不足所带来的挑战，提高分类的准确性。考虑线性映射 $\mathbf{W} \in R^{D \times d}$ ，其中 D 和 d 分别是变换之前和之后数据的维度。新损失函数的优点在于，通过在约束条件 $\mathbf{W}^T \mathbf{W} = \mathbf{I}$ 下优化目标，能够避免了对类内散度矩阵 $\mathbf{S}_w$ 求逆的需求，从而有效规避了小样本量问题可能带来的影响。为了解决上述优化问题，我们可以引入拉格朗日函数：

$$\mathcal{L} = \text{tr}(\mathbf{W}^T (\mathbf{S}_b - \mathbf{S}_w) \mathbf{W}) + \lambda (\text{tr}(\mathbf{W}^T \mathbf{W}) - m) \quad (15)$$

两边同时对 $\mathbf{W}$ 求导，

$$(\mathbf{S}_b - \mathbf{S}_w)\mathbf{W} = \lambda\mathbf{W}, \quad (16)$$

这就意味着 λ 是 $\mathbf{S}_b - \mathbf{S}_w$ 的特征值，而 $\mathbf{W}$ 是对应的特征向量。因此，当 $\mathbf{W}$ 由前 d 个最大特征向量组成时，损失函数 J 达到最大值。在这里，避免了直接计算 $\mathbf{S}_w$ 的逆矩阵，从而减轻小样本问题的影响。

相比传统 LDA 方法，Euler-LDA 在捕捉非线性信息、增强类别区分度和解决小样本问题方面具有显著优势。1) 欧拉表示通过将数据映射到复空间，更准确地刻画了数据中的非线性关系，显著提升了算法处理非线性分布数据的能力。2) 欧拉表示接近理想的鲁棒核，能够扩大类内和类间的度量差异，同时提供比欧氏距离更稳定且准确的相似性度量，从而增强了类别区分度。3) Euler-LDA 引入了新的目标损失函数，通过最大化类间散布与类内散布的差异，有效避免了矩阵奇异性问题，显著缓解了小样本问题对算法性能的影响。

5. 实验

5.1. 实验数据集

实验在 Coil-100、ORL 和 AR 三个数据库上进行，具体信息如下：

Coil-100 数据集：这是一个包含多个物体的彩色图像数据集，物体以 360 度不同角度拍摄，形成了 7200 张图像。每张图像的分辨率为 128×128 像素，适合用于物体识别和角度变化研究。

ORL 人脸数据集：该数据集由英国剑桥大学的 Olivetti 研究实验室于 1992 年发布，包含 40 个个体的 400 张人脸图像。这些图像在不同的时间、光照下拍摄，并涵盖了多种面部表情和环境变化。每张图像的大小为 32×32 像素，适合进行小样本学习和面部识别实验。

AR 人脸数据集：该数据集包含 126 位个体的人脸图像，其中男性 70 人，女性 56 人。图像呈现了多种变化，如不同的面部表情、光照条件以及遮挡(例如佩戴墨镜或围巾)。总计收录了超过 4000 张彩色图像，展现了高度的多样性和挑战性。

这些数据库涵盖了不同的人脸表情和物体的角度变化，为算法的评估提供了丰富且多样化的测试数据，数据库中的一些图片展示在图 1 中。



Figure 1. From top to bottom are some samples from the Coil-100, ORL, and AR datasets

图 1. 从上到下是 Coil-100、ORL 和 AR 数据集的一些样本

5.2. 参数设置

在每个数据集中，通过随机挑选每个类别中的 p 张图像来构建训练集，其余图像作为测试集。其中 p 值代表训练样本的图像数量。对于给定的 p ，降维的子空间维度从 10 增加到 100，步长为 10。然后，对于每个子空间维度，我们计算相应的识别率。此过程可以看成是一个重复的循环。对于训练样本数 p ，我们计算 10 个周期。这样每个子空间维度就有 10 个识别率，然后我们取它们的平均值作为当前 p 和子空间维度的识别率。最后，我们从最佳子空间维度中取最佳识别率作为训练样本的结果。

对于每个数据集的各个类别，我们随机挑选 p 张图像作为训练样本，其余图像用于测试。在子空间维度范围从 10 到 100 (间隔为 10) 内，计算各维度下的识别率。将该过程重复 10 次，并对识别率取平均识别率。最终，以最高平均识别率对应的最佳子空间维度，作为当前训练样本数下的实验结果。此外，实验中我们使用 K 最近邻(KNN)分类器进行分类，并且新方法和本文对比方法的 K 为 1。

5.3. 分类实验结果

我们在 AR、ORL 和 Coil-100 人脸数据集，将 Euler-LDA 与 LDA、KDA、LDAMMC [12]、NLDA [13]、TSLDA [14]、和 ALDE [15] 进行比较。表 1 至表 3 列出了上述方法在三个数据集上的最佳识别率、标准差和最佳子空间维度。对于每个类别，AR 数据集的 p 值为 4、3 和 5；ORL 数据集 p 值为 3、4 和 5；Coil-100 数据集的 8、10 和 12。

根据识别率的分析，随着训练样本数量 p 的增加，大多数算法在 AR、ORL 和 Coil-100 数据集上的表现呈现出一致的提升趋势，这反映了更多的训练样本有助于模型学习到更丰富的特征，从而提高了识

Table 1. Recognition accuracy and optimal dimensionality of the AR database

表 1. AR 数据库的识别准确率和最优维度

方法	$p = 3$	$p = 4$	$p = 5$
LDA	85.03 ± 2.05 (10)	89.97 ± 1.35 (10)	92.62 ± 0.43 (10)
KDA	87.32 ± 0.89 (10)	91.47 ± 0.18 (10)	93.80 ± 0.32 (10)
LDAMMC	83.18 ± 0.75 (100)	91.05 ± 0.48 (100)	94.75 ± 0.62 (100)
NLDA	87.02 ± 0.76 (100)	91.94 ± 0.49 (100)	94.60 ± 0.42 (100)
TSLDA	87.88 ± 0.35 (100)	91.61 ± 0.33 (100)	93.37 ± 0.27 (100)
ALDE	86.44 ± 0.59 (100)	92.80 ± 0.80 (100)	95.34 ± 0.79 (90)
Euler-LDA	94.14 ± 1.52 (90)	94.59 ± 1.35 (90)	95.97 ± 1.91 (90)

Table 2. Recognition accuracy and optimal dimensionality of the ORL database

表 2. ORL 数据库的识别准确率和最优维度

方法	$p = 3$	$p = 4$	$p = 5$
LDA	83.33 ± 1.76 (10)	89.17 ± 0.84 (10)	93.00 ± 1.32 (10)
KDA	84.17 ± 2.03 (20)	88.75 ± 0.72 (10)	93.67 ± 1.04 (20)
LDAMMC	83.57 ± 1.89 (40)	88.89 ± 2.30 (30)	93.17 ± 1.60 (40)
NLDA	87.02 ± 1.61 (40)	92.78 ± 0.89 (40)	96.33 ± 1.04 (40)
TSLDA	84.05 ± 4.64 (60)	88.75 ± 1.10 (30)	92.33 ± 0.76 (30)
ALDE	84.29 ± 2.34 (40)	89.31 ± 1.93 (40)	92.83 ± 1.89 (40)
Euler-LDA	90.35 ± 3.55 (50)	93.61 ± 1.81 (70)	96.16 ± 1.46 (50)

Table 3. Recognition accuracy and optimal dimensionality of the Coil-100 database**表 3.** Coil-100 数据库的识别准确率和最优维度

方法	$p = 8$	$p = 10$	$p = 12$
LDA	61.09 ± 0.69 (10)	65.49 ± 0.16 (10)	68.09 ± 0.82 (10)
KDA	74.30 ± 0.70 (10)	79.78 ± 0.66 (10)	81.99 ± 1.43 (10)
LDAMMC	85.92 ± 0.53 (30)	89.19 ± 0.21 (40)	90.40 ± 0.87 (100)
NLDA	71.33 ± 0.60 (20)	71.53 ± 0.47 (30)	71.40 ± 0.60 (40)
TSLDA	63.87 ± 0.46 (40)	65.73 ± 0.07 (100)	65.89 ± 0.62 (50)
ALDE	85.23 ± 0.39 (20)	88.67 ± 0.20 (30)	90.10 ± 1.08 (40)
Euler-LDA	89.36 ± 1.87 (40)	90.61 ± 0.83 (30)	92.93 ± 2.31 (40)

别能力。同时，这也表明，随着样本数量的增加，模型能够更好地捕捉到数据中的潜在规律，减少过拟合的风险，进而提升分类性能。在 AR 和 ORL 数据集中，当训练样本数 p 分别为 4 和 10 时，算法分别达到了最高识别率 95.97% 和 96.16%；而在 Coil-100 数据集中，当训练样本数 p 为 12 时，最高识别率为 92.93%。实验结果表明，Euler-LDA 在所有参数设置下的识别率均优于传统 LDA 及其改进方法。

此外，我们还针对子空间维度的变化评估了这些方法的性能。具体而言，在给定的训练样本 p 固定的情况下，将子空间维度设置为 10 到 100 (步长为 10)，并记录不同维度下的识别率。实验结果显示了各维度对应的识别率变化趋势，其具体曲线如图 2 到图 4 所示。

Euler-LDA 的识别率远超过原始 LDA 方法及其变体，显示了其在复杂情况下的有效性和可靠性。在 AR、ORL 和 Coil-100 数据集上，当训练样本数量 p 分别设置为 4、3 和 12 时，Euler-LDA 的识别率相较于原始 LDA 分别提高了 4.62%、7.02% 和 24.84%。这可能与 Euler-LDA 能够更好地捕捉数据中复杂的非线性结构有关。

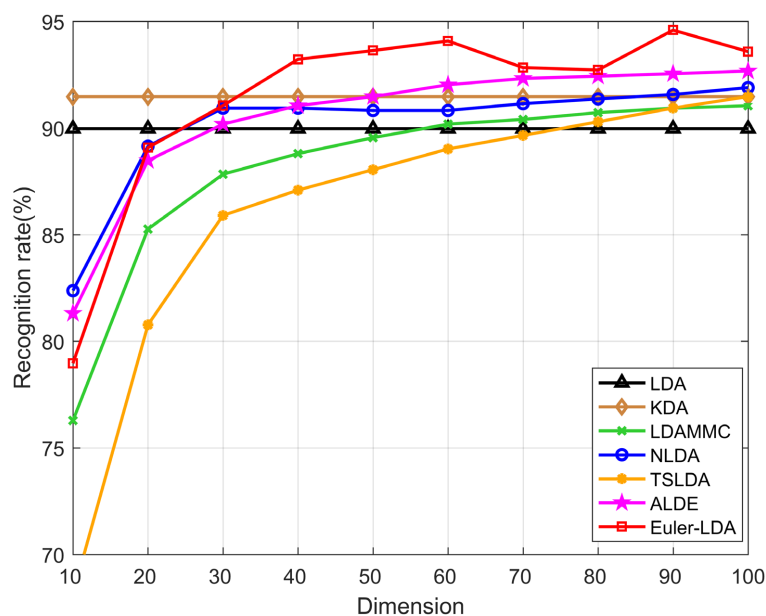


Figure 2. Comparison of recognition rates on AR dataset (with training samples $p = 4$)

图 2. AR 数据集(训练样本 $p = 4$)上的识别率比较

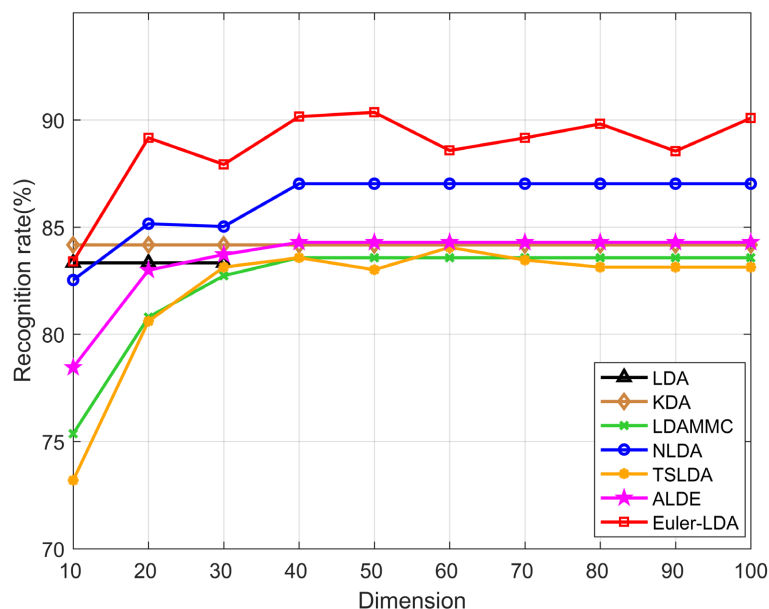


Figure 3. Comparison of recognition rates on ORL dataset (with training samples $p = 4$)

图 3. ORL 数据集(训练样本 $p = 4$)上的识别率比较

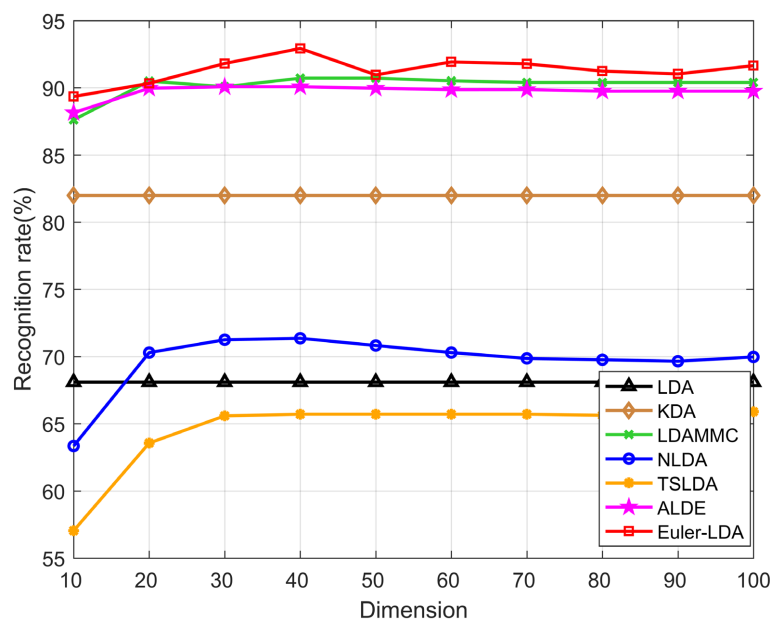


Figure 4. Comparison of recognition rates on ORL dataset (with training samples $p = 4$)

图 4. ORL 数据集(训练样本 $p = 4$)上的识别率比较

5.4. 数据可视化

为直观展示 LDA 和 Euler-LDA 在分类性能上的差异,我们在 Coil-100 数据集上进行了数据可视化实验,将降维后的数据映射到二维空间。实验中,训练样本数 p 设置为 12,测试样本数为 60,选取数据集中 5 个类别,每个类别包含 60 个测试样本点。结果如图 5 所示。通过对比可以发现, Euler-LDA 在区分不同类别方面比传统 LDA 表现更优,具体表现为不同类别的数据点分布更为分散,而同类别的数据点

聚集性更高。

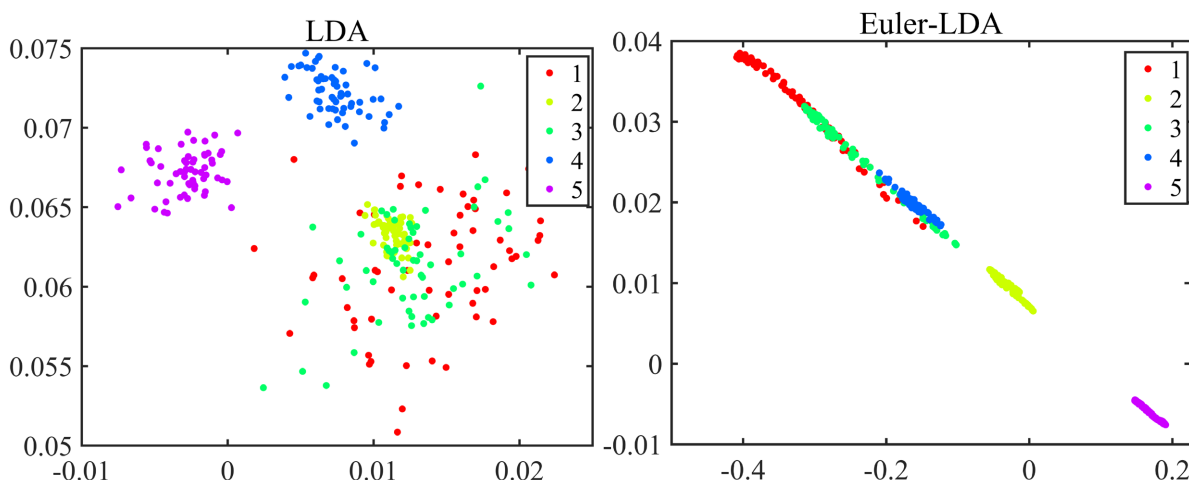


Figure 5. The visualizations experiments for LDA and Euler-LDA, in sequence

图 5. 依次是 LDA 和 Euler-LDA 的可视化实验

从 LDA 的聚类图可以看出,不同颜色的点群显示出一定程度的可分性,部分类别之间的边界较为清晰,但仍存在一定的重叠区域。这表明在实际应用中,某些样本可能无法完全正确分类。总体而言, LDA 能够在二维空间内有效区分不同类别的点群,部分类别之间的间隔较大,显示出良好的分类性能。然而,某些类别的重叠区域可能导致分类精度的下降。

相比之下, Euler-LDA 结合欧拉表示的优势,将数据映射到复空间,通过复数的几何特性更准确地描述数据的非线性关系,显著增强了对数据特征的提取能力。其聚类效果在视觉上更具层次性,点群分布更加稀疏且紧凑,使得类别之间的边界更加清晰。这种紧凑的分布模式表现出更强的可分性,从而有效提升了分类性能和鲁棒性。

6. 总结

线性判别分析(LDA)在提取数据判别信息方面表现出色。然而,传统基于欧氏距离的 LDA 难以有效捕捉非线性分布数据的复杂特征。为解决传统线性判别分析(LDA)难以捕捉非线性分布数据特征的问题,本章提出了一种近似于核方法的新方法,称为基于欧拉表示的线性判别分析模型(Euler-LDA)。Euler-LDA 结合欧拉表示的优势,将数据映射到复空间,通过复数的几何特性更准确地描述数据的非线性关系,从而增强了算法处理非线性分布数据的能力。同时,欧拉表示能够逼近理想的鲁棒核,显著扩大类内与类间的度量差异,提供比欧氏距离更稳定且更精确的相似性度量,进一步提升类别区分能力。该方法在保留传统 LDA 处理线性数据优势的基础上,显著增强了其对非线性分布数据的适应性,为复杂数据的降维与分类问题提供了一种更稳健且高效的解决方案。

综上所述,相比传统 LDA 方法, Euler-LDA 在三个方面具有显著优势: 1) 捕捉非线性判别信息: 欧拉表示能够更准确地描述数据中的非线性关系,从而提高算法处理非线性分布数据的能力。2) 增强类别区分度: 欧拉表示近似于理想的鲁棒核,不仅能扩大类内和类间的度量差异,还能提供比欧氏距离更稳定和准确的相似性度量。3) Euler-LDA 引入了新的目标损失函数,通过最大化类间散布与类内散布的差异,有效避免了矩阵奇异性问题,显著缓解了小样本问题对算法性能的影响。

本文通过实验对 Euler-LDA 的性能进行了全面评估。在多个数据集上的对比实验中, Euler-LDA 展现了出色的识别精度。此外,通过可视化分析,进一步验证了模型的鲁棒性和稳定性。

参考文献

- [1] Van Der Maaten, L., Postma, E.O. and Van Den Herik, H.J. (2009) Dimensionality Reduction: A Comparative Review. *Journal of Machine Learning Research*, **10**, 66-71.
- [2] Abdi, H. and Williams, L.J. (2010) Principal Component Analysis. *WIREs Computational Statistics*, **2**, 433-459. <https://doi.org/10.1002/wics.101>
- [3] He, X. and Niyogi, P. (2003) Locality Preserving Projections. *Proceedings of the 17th International Conference on Neural Information Processing Systems*, Whistler, 9-11 December 2003, 153-160.
- [4] Balakrishnama, S. and Ganapathiraju, A. (1998) Linear Discriminant Analysis—A Brief Tutorial. Institute for Signal and Information Processing, 1-8.
- [5] Schölkopf, B., Smola, A. and Müller, K. (1997) Kernel Principal Component Analysis. In: Gerstner, W., et al., Eds., *Artificial Neural Networks*, Springer, 583-588. <https://doi.org/10.1007/bfb0020217>
- [6] Pekalska, E. and Haasdonk, B. (2009) Kernel Discriminant Analysis for Positive Definite and Indefinite Kernels. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **31**, 1017-1032. <https://doi.org/10.1109/tpami.2008.290>
- [7] Fitch, A.J., Kadyrov, A., Christmas, W.J. and Kittler, J. (2005) Fast Robust Correlation. *IEEE Transactions on Image Processing*, **14**, 1063-1073. <https://doi.org/10.1109/tip.2005.849767>
- [8] Huber, P.J. (1981) Wiley Series in Probability and Mathematics Statistics. In: Huber, P.J., Ed., *Robust Statistics*, John Wiley & Sons, Inc., 309-312.
- [9] Liu, Y., Gao, Q., Han, J. and Wang, S. (2018) Euler Sparse Representation for Image Classification. *Proceedings of the AAAI Conference on Artificial Intelligence*, **32**, 3691-3697. <https://doi.org/10.1609/aaai.v32i1.11670>
- [10] Liwicki, S., Tzimiropoulos, G., Zafeiriou, S. and Pantic, M. (2012) Euler Principal Component Analysis. *International Journal of Computer Vision*, **101**, 498-518. <https://doi.org/10.1007/s11263-012-0558-z>
- [11] Long, T., Sun, Y., Gao, J., Hu, Y. and Yin, B. (2020) Locality Preserving Projection Based on Euler Representation. *Journal of Visual Communication and Image Representation*, **70**, Article ID: 102796. <https://doi.org/10.1016/j.jvcir.2020.102796>
- [12] Li, H., Jiang, T. and Zhang, K. (2003) Efficient and Robust Feature Extraction by Maximum Margin Criterion. *IEEE Transactions on Neural Networks*, **17**, 157-165.
- [13] Lu, G. and Wang, Y. (2012) Feature Extraction Using a Fast Null Space Based Linear Discriminant Analysis Algorithm. *Information Sciences*, **193**, 72-80. <https://doi.org/10.1016/j.ins.2012.01.015>
- [14] Chen, Y., Tao, X., Xiong, C., et al. (2018) An Improved Method of Two Stage Linear Discriminant Analysis. *KSII Transactions on Internet and Information Systems (TIIS)*, **12**, 1243-1263.
- [15] Smith, R.S., Kittler, J., Hamouz, M. and Illingworth, J. (2006) Face Recognition Using Angular LDA and SVM Ensembles. *18th International Conference on Pattern Recognition (ICPR'06)*, Vol. 3, 1008-1012. <https://doi.org/10.1109/icpr.2006.529>