

Geometric interpretation of Fisher's linear discriminant analysis through communication theory

Jun FUJIKI and Masaru TANAKA

Fukuoka University

8-19-1, Nanakuma, Jonan-ku, Fukuoka 814-0180

fujiki@fukuoka-u.ac.jp, sieger@math.sci.fukuoka-u.ac.jp

Hitoshi SAKANO

NTT Data

3-3-9, Toyosu, Koutou-ku, Tokyo 135-8671

sakanoh@nttdata.co.jp

Akisato KIMURA

NTT Communication Science Laboratories

3-1, Morinosato Wakamiya Atsugi-shi, Kanagawa 243-0198

akisato@ieee.org

Abstract

This paper provides a geometrical aspect of Fisher's linear discriminant analysis (FLDA), which has been widely used owing to its simple formulation and low computational costs. Our approach is based on a new framework of pattern recognition that can be modeled by a communication of class information. This model is quite different from a commonly used framework of pattern recognition as a mapping from the set of patterns to the set of classes. In the new framework, patterns can be regarded as class information with redundant encoding. We show that the geometry of two class FLDA can be described via communication theory of noisy channel.

1 Introduction

Fisher's linear discriminant analysis (FLDA) [7] has been widely used as a discriminative feature extractor in the fields of pattern recognition, computer vision and machine learning [3, 10] for a long time owing to its simple formulation and low computational costs. Lots of extensions and modifications of FLDA have ever been proposed [2, 5, 8, 9, 11, 12, 13, 14, 15], and these FLDA-based methods suggest that FLDA is a fundamental and important method in pattern recognition. Therefore, to understand geometries of pattern recognition, understanding a geometry of FLDA is very important. Then, this paper provides a geometrical aspect of two class FLDA by a new framework of pattern recognition.

Our new framework of pattern recognition is inspired by the communication theory [1, 4]: We regard pattern recognition as a communication of class information on a noisy channel, in contrast to a common framework of pattern recognition as a mapping from a pattern space consisting of spatial/temporal data such as images and sounds, into class space. In the framework, pattern recognition is regarded as a mapping from a pattern into what a pattern represents that belongs to some discrete class. Usually, the dimension of pattern space is very large, mean while the dimension of class space is small. Therefore, pattern recognition as

a compressive mapping from redundant pattern space into concise class space can be viewed as dimensional reduction and/or information compression. In the context of pattern recognition, the mapping from pattern space to class space should be carefully constructed to achieve high recognition rate. This mapping is generally decomposed of two steps: one is from a pattern space into a feature space, and the other from a feature space into a class space, which is called "feature extraction" in general.

This paper gives a new framework of pattern recognition from the point of view of communication theory. In a common, framework of pattern recognition, a pattern space is the basis of theoretical analysis, meanwhile in our framework, a class space plays a central role in the analysis. Every class as an element of a class space is given a priori, and a pattern can be regarded as class information with redundant encoding. Pattern recognition is a process of decoding observed patterns with channel noises into class labels by projecting it to small dimensional space (Fig. 1). The image of a pattern should have discrete values, however, noise contamination makes it continuously. Then, each label is embedded to a continuous space and the value of continuous label is regarded as a label likelihood. This leads to a new framework of pattern recognition, such that pattern recognition can be regarded as communication of the class label, then, pattern recognition can be understood in the framework of the communication theory.

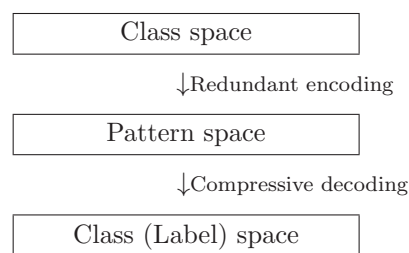


Figure 1. Proposed pattern recognition process.

Making use of the above idea, we can understand two class FLDA from the point of the view of information geometry.

2 Fisher's linear discriminant analysis

In this section, we explain FLDA, which is a linear method of supervised learning. FLDA projects data into a small dimensional space by emphasizing the separability of classes, that is, the data belonging to the same class would be located near and the data belonging to the different classes would be located far.

Assuming that N data are classified to C classes. Let n_c be the number of data classified to the class c , and they are named $\mathbf{x}_1^{(c)}, \dots, \mathbf{x}_{n_c}^{(c)}$ ($N = \sum_{c=1}^C n_c$). Then, the mean and the variance-covariance matrix of the class c be

$$\boldsymbol{\mu}^{(c)} = \frac{1}{n_c} \sum_{i=1}^{n_c} \mathbf{x}_i^{(c)},$$

$$\Sigma_W^{(c)} = \frac{1}{n_c} \sum_{i=1}^{n_c} \left(\mathbf{x}_i^{(c)} - \boldsymbol{\mu}^{(c)} \right) \left(\mathbf{x}_i^{(c)} - \boldsymbol{\mu}^{(c)} \right)^\top,$$

respectively. Here, $\boldsymbol{\mu}^{(c)}$ and $\Sigma_W^{(c)}$ are called a *class mean* and a *within class variance*, respectively.

Let the mean $\boldsymbol{\mu}$ and the variance-covariance matrix Σ_T of all data be

$$\boldsymbol{\mu} = \frac{1}{N} \sum_{c=1}^C \sum_{i=1}^{n_c} \mathbf{x}_i^{(c)} = \frac{1}{N} \sum_{c=1}^C n_c \boldsymbol{\mu}^{(c)},$$

$$\Sigma_T = \frac{1}{N} \sum_{c=1}^C \sum_{i=1}^{n_c} \left(\mathbf{x}_i^{(c)} - \boldsymbol{\mu} \right) \left(\mathbf{x}_i^{(c)} - \boldsymbol{\mu} \right)^\top,$$

respectively. Let Σ_W , called a *within-class variance*, be the average of the variance-covariance matrix of every classes, and let Σ_B , called a *between-class variance*, be the variance-covariance matrix when all data is assumed to concentrate to the mean of a corresponding class. These are represented by

$$\Sigma_W = \frac{1}{N} \sum_{c=1}^C n_c \Sigma_W^{(c)}, \quad (1)$$

$$\Sigma_B = \frac{1}{N} \sum_{c=1}^C n_c \left(\boldsymbol{\mu}^{(c)} - \boldsymbol{\mu} \right) \left(\boldsymbol{\mu}^{(c)} - \boldsymbol{\mu} \right)^\top,$$

respectively. Then, there holds $\Sigma_T = \Sigma_W + \Sigma_B$.

We consider two class FLDA. In this case, data are projected (compressed) to a 1-dimensional linear space spanned by $\mathbf{w} \neq \mathbf{0}$. For projected data, there holds

$$\mathbf{w}^\top \Sigma_T \mathbf{w} = \mathbf{w}^\top \Sigma_W \mathbf{w} + \mathbf{w}^\top \Sigma_B \mathbf{w}.$$

In the discrimination, a between-class variance is more important than a within-class variance. Then, $\mathbf{w}$ is chosen to maximize the ratio $\frac{\mathbf{w}^\top \Sigma_B \mathbf{w}}{\mathbf{w}^\top \Sigma_W \mathbf{w}}$, and $\mathbf{w}$ is explicitly given in the form

$$\mathbf{w} = \Sigma_W^{-1} \left(\boldsymbol{\mu}^{(1)} - \boldsymbol{\mu}^{(2)} \right), \quad (2)$$

by solving a general eigenvalue problem as the eigenvector corresponding to the largest eigenvalue of the matrix $\Sigma_W^{-1} \Sigma_B$.

3 Overview of communication theory

In this section, we overview the communication theory [1, 4].

3.1 Entropy and transinformation

Let $\mathbf{X}$ be a random variable of a signal $\mathbf{x}$ generated by a probability density function (pdf) $p_{\mathbf{X}}(\mathbf{x})$. The ambiguity of signal is measured by the entropy of $\mathbf{X}$ as $H(\mathbf{X}) = - \int p_{\mathbf{X}}(\mathbf{x}) \log p_{\mathbf{X}}(\mathbf{x}) d\mathbf{x}$.

For two random variables $\mathbf{X}$ and $\mathbf{Y}$ corresponding to signals $\mathbf{x}$ and $\mathbf{y}$, respectively, the joint entropy of the pair $\mathbf{X}$ and $\mathbf{Y}$ and the conditional entropy of $\mathbf{Y}$ given $\mathbf{X}$ are defined by

$H(\mathbf{X}, \mathbf{Y}) = - \int p_{\mathbf{X}, \mathbf{Y}}(\mathbf{x}, \mathbf{y}) \log p_{\mathbf{X}, \mathbf{Y}}(\mathbf{x}, \mathbf{y}) d\mathbf{x} d\mathbf{y}$, and $H(\mathbf{Y}|\mathbf{X}) = - \int p_{\mathbf{X}}(\mathbf{x}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) \log p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) d\mathbf{x} d\mathbf{y}$, respectively, where the pair $\mathbf{x}$ and $\mathbf{y}$ are generated by the pdf $p_{\mathbf{X}, \mathbf{Y}}(\mathbf{x}, \mathbf{y})$, and $p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) = \frac{p_{\mathbf{X}, \mathbf{Y}}(\mathbf{x}, \mathbf{y})}{p_{\mathbf{X}}(\mathbf{x})}$

be the conditional pdf of $\mathbf{y}$ given $\mathbf{x}$. Let $I(\mathbf{X}, \mathbf{Y})$, called a *mutual information* of $\mathbf{X}$ and $\mathbf{Y}$, be defined by the difference of entropies (signal ambiguity) as $H(\mathbf{Y}) - H(\mathbf{Y}|\mathbf{X}) = H(\mathbf{X}) - H(\mathbf{X}|\mathbf{Y})$. The mutual information gives the common information between $\mathbf{X}$ and $\mathbf{Y}$, and it equals 0 when $\mathbf{X}$ and $\mathbf{Y}$ are independent one another.

Let a signal $\mathbf{y}$ be a transmitted signal of $\mathbf{x}$ contaminated by a noise $\mathbf{n}$ through continuous channel (the channel transmits continuous signals), that is, there holds $\mathbf{y} = \mathbf{x} + \mathbf{n}$. Because of the noise, $\mathbf{x}$ cannot be exactly restored by $\mathbf{y}$. Let $\mathbf{X}$ and $\mathbf{N}$ be random variables of a signal $\mathbf{x}$ and a noise $\mathbf{n}$, respectively. We assume $\mathbf{X}$ and $\mathbf{N}$ are independent. The loss of information due to noise is $H(\mathbf{X}|\mathbf{Y}) = H(\mathbf{N})$, then, transmitted information of $\mathbf{X}$ equals to the mutual information $J = I(\mathbf{X}, \mathbf{Y})$. Then, J is also called *transinformation* from the point of view of data transmission.

3.2 Riemannian signal space

Let a noise vector $\mathbf{n}$ be an additive Gaussian noise of mean $\mathbf{0}$ depends on the signal $\mathbf{x}$, that is, the conditional pdf of $\mathbf{n}$ given $\mathbf{x}$ be

$$p_{\mathbf{N}|\mathbf{X}}(\mathbf{n}|\mathbf{x}) = A(\mathbf{x}) \exp \left\{ -\frac{1}{2} \mathbf{n}^\top V^{-1}(\mathbf{x}) \mathbf{n} \right\} \quad (3)$$

be holds, where $V(\mathbf{x})$ be the variance-covariance matrix of $\mathbf{n}$ given $\mathbf{x}$, and $A(\mathbf{x}) = (2\pi)^{-n/2} (\det V(\mathbf{x}))^{-1/2}$ be a normalization factor.

By the noise, signals are shifted, and the way of shift of signals can be represented by the ellipsoid $\mathbf{n}^\top V^{-1}(\mathbf{x}) \mathbf{n} = \text{const.}$, which corresponds to Gaussian noise as large (small) noise is incident to the direction along long (short) axis. The metric in the signal space is defined by the characteristic of the signal shift as the square of an infinitesimal distance between two signals $\mathbf{x}$ and $\mathbf{x}' = \mathbf{x} + d\mathbf{x}$ is represented by $ds^2 = \{d(\mathbf{x}, \mathbf{x}')\}^2 = d\mathbf{x}^\top G(\mathbf{x}) d\mathbf{x}$, where $G(\mathbf{x}) = \frac{1}{n} V^{-1}(\mathbf{x})$ and n is the dimension of the signal space. The matrix $G(\mathbf{x})$ is called a *Gramian matrix*. It is called a Riemannian space that the space where the Gramian matrix is defined at all points. When the signal space is a Riemannian space, the space is called a *Riemannian signal space*.

By the Gramian matrix, the square-length of the noise vector $\mathbf{n}$ is $\mathbf{n}^\top G(\mathbf{x})\mathbf{n}$. Therefore, the average of the square-length is $\overline{\mathbf{n}^\top G(\mathbf{x})\mathbf{n}} = \text{tr}[G(\mathbf{x})V(\mathbf{x})] = \frac{1}{n}\text{tr}[\mathbf{I}] = 1$, where $\mathbf{I}$ is the identity matrix. Then, the metric is defined so that the average of the noise-length equals to 1. When the dimension of the signal space is sufficiently large, the noise-length is 1 for almost all noise because of low of large number. That is, almost all signals are shifted to be on the unit hypersphere of center $\mathbf{x}$. The unit hypersphere is called a *noise hypersphere*. Note that the shape of a noise hypersphere is depend on the signal $\mathbf{x}$.

When Gaussian noise vector $\mathbf{n}$ is sufficiently small, by neglecting higher order term, $H(\mathbf{X}|\mathbf{Y})$ is approximated by $H(\mathbf{N}|\mathbf{X})$ [1], transinformation is approximated by $J = H(\mathbf{X}) - H(\mathbf{N}|\mathbf{X})$.

Let $g(\mathbf{x})$ be $\det G(\mathbf{x})$, there holds

$$p_{\mathbf{N}|\mathbf{X}}(\mathbf{n}|\mathbf{x}) = \left(\frac{n}{2\pi}\right)^{\frac{n}{2}} \sqrt{g(\mathbf{x})} \exp\left\{-\frac{n}{2}\mathbf{n}^\top G(\mathbf{x})\mathbf{n}\right\},$$

and therefore the entropy of noise at $\mathbf{x}$ in signal space is computed by using Eq. (3) as

$$H(\mathbf{N}|\mathbf{x}) = -\log\left(b_n \sqrt{g(\mathbf{x})}\right)$$

where $b_n = \left(\frac{n}{2\pi e}\right)^{\frac{n}{2}}$. Then, the average transinformation of the signal space is

$$J = H(\mathbf{X}) - H(\mathbf{N}|\mathbf{X}) = \int p_{\mathbf{X}}(\mathbf{x}) \log \frac{b_n \sqrt{g(\mathbf{x})}}{p_{\mathbf{X}}(\mathbf{x})} d\mathbf{x}.$$

Because the average transinformation depends on $p_{\mathbf{X}}(\mathbf{x})$, we control $p_{\mathbf{X}}(\mathbf{x})$ to maximize the average transinformation in order to realize the best signal transmission. The maximized average transinformation satisfies the condition $I = \max_{p_{\mathbf{X}}(\mathbf{x})} J$, and the value I is called a *channel capacity* of the signal space. By Lagrange's multiplier method, the channel capacity can be exactly obtained as

$$I = \log(b_n U)$$

where $U = \int \sqrt{g(\mathbf{x})} d\mathbf{x}$ at $p_{\mathbf{X}}(\mathbf{x}) = \frac{1}{U} \sqrt{g(\mathbf{x})}$. Because the volume of the noise hypersphere is $\frac{\pi^{\frac{n}{2}}}{\sqrt{g(\mathbf{x})}}$, $\sqrt{g(\mathbf{x})}$ is inversely proportional to the volume of the noise hypersphere, that is, proportional to the number of the noise hypersphere packed in the neighborhood of $\mathbf{x}$. U is the average of $\sqrt{g(\mathbf{x})}$ over the signal space. Therefore, $b_n U$ represents the number of the noise hypersphere packed in the signal space, and then, the same number of signals can be correctly discriminated. Consequently, the channel capacity is equivalent to the logarithm of the number of the noise hypersphere.

3.3 Imbedding/reduction mapping

In this subsection, we explain a continuous channel as a sequential mapping in the signal space (Fig. 2). At First, signals in a Riemannian signal space S^n is redundantly encoded as elements in a high dimensional channel signal space S^m . The elements in S^m are contaminated by noises. After that, noise-contaminated signals in S^m are decoded as elements in S^n . For the

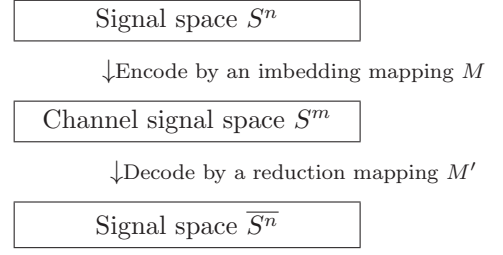


Figure 2. Channel model.

sake of clarity, we note $\overline{S^n}$ as the Riemannian signal space consists of decoded elements.

To utilize the channel efficiently, the best imbedded mapping $M : S^n \rightarrow S^m$, and the best reduction mapping $M' : S^m \rightarrow \overline{S^n}$, should be chosen carefully.

Let $\mathbf{f}$ be a mapping from $\mathbf{x} \in S^n$, into $\mathbf{y} \in S^m$. The mapping $\mathbf{y} = \mathbf{f}(\mathbf{x})$ is called an *imbedding mapping*. The mapping $\mathbf{f}(\mathbf{x})$ is locally linearized by $d\mathbf{y} = T(\mathbf{x})d\mathbf{x}$, where $T(\mathbf{x}) = \frac{\partial \mathbf{f}}{\partial \mathbf{x}}$. Here, T is called a *tangent matrix* because the set of T 's columns is the basis of the tangent space of $\mathbf{y} = \mathbf{f}(\mathbf{x})$.

On the other hand, the signal $\mathbf{y}$ is inversely transformed to signal $\mathbf{x} \in \overline{S^n}$. Let $\mathbf{h}$ be a mapping from $\mathbf{y}$ into $\mathbf{x}$. The mapping $\mathbf{x} = \mathbf{h}(\mathbf{y})$ is called a *reduction mapping*. The mapping $\mathbf{h}(\mathbf{x})$ is locally linearized by $d\mathbf{x} = R(\mathbf{y})d\mathbf{y}$, where $R(\mathbf{y}) = \frac{\partial \mathbf{h}}{\partial \mathbf{y}}$. Here, R is called a *reduction matrix*. Note that $R(\mathbf{x})T(\mathbf{x}) = \mathbf{I}$ holds because the decoded signals should be the same as the original.

Let a noise $\mathbf{m}$ in S^m be sufficiently small, then $\mathbf{h}(\mathbf{y} + \mathbf{m}) = \mathbf{x} + R(\mathbf{x})\mathbf{m}$ holds. Then, the noise in $\overline{S^n}$ is represented by $\mathbf{n} = R(\mathbf{x})\mathbf{m}$. Then, if the noise satisfies $R(\mathbf{x})\mathbf{m} = \mathbf{0}$, $\mathbf{x}$ is correctly decoded by the reduction mapping. Note that the set of $\mathbf{m}$ constructs the null space of $R(\mathbf{x})$, and the reduction mapping $\mathbf{h}$ locally projects $\mathbf{y}$ into $\overline{S^n}$ along the null space of $R(\mathbf{x})$.

3.4 Optimal reduction mapping

Let a noise $\mathbf{m}$ in S^m be an additive Gaussian noise of a mean $\mathbf{0}$ and a variance-covariance matrix $V(\mathbf{y})$.

By introducing a Gramian matrix at $\mathbf{y}$ in S^m as $F(\mathbf{y}) = \frac{1}{m}V^{-1}(\mathbf{y})$, S^m can be regarded as a Riemannian signal space. The Gramian matrix in S^m provides the metric in $\overline{S^n}$ as $G = \frac{m}{n}(RF^{-1}R^\top)^{-1}$, since the mapping from S^m into $\overline{S^n}$ is the reduction mapping R . To realize an efficient transmission, the channel capacity $I = \log\{b_n \int \sqrt{g(\mathbf{x})} d\mathbf{x}\}$ should be maximized, that is, $g(\mathbf{x})$ should be maximized.

The reduction mapping which maximize $g(\mathbf{x})$ is called an *optimal reduction mapping* and the decoding by the optimal reduction mapping is called an *optimal signal extraction method*. By Lagrange's multiplier method, The optimal reduction mapping is derived as

$$\tilde{R} = \frac{m}{n} \tilde{G}^{-1} T^\top F.$$

where $\tilde{G} = \frac{m}{n} T^\top F T \left(= \frac{m}{n} \left(\tilde{R} F^{-1} \tilde{R}^\top \right)^{-1} \right)$, and this is the pdf of noise in $\overline{S^n}$ under the optimal reduction mapping. The $\tilde{G}$ is called an *optimal Gramian matrix*.

It is known that the optimal Gramian matrix is proportional to Fisher information matrix for estimating $\mathbf{x}$ from $\mathbf{y}$.

4 Geometry of two class FLDA

In this section, we explain the geometry of two class FLDA, which is equivalent to the optimal reduction mapping in the context of communication theory.

Let patterns be probabilistically generated from a pdf $p(\mathbf{x}; s)$, which is a mixture of two pdfs $p_1(\mathbf{x})$ and $p_2(\mathbf{x})$ at a rate of $\frac{1-s}{2} : \frac{1+s}{2}$ as $p(\mathbf{x}; s) = \frac{1-s}{2}p_1(\mathbf{x}) + \frac{1+s}{2}p_2(\mathbf{x})$.

Let $p_i(\mathbf{x})$'s are assumed to be a Gaussian distribution of mean $\boldsymbol{\mu}^{(i)}$, the mixture $p(\mathbf{x}; s)$ is also a Gaussian distribution of mean $\boldsymbol{\mu} = \frac{\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2}{2} + s\frac{\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1}{2}$. Then, the mixture rate s can be computed by $s = \frac{\langle 2\boldsymbol{\mu} - (\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2) | \boldsymbol{\mu} \rangle}{\langle \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1 | \boldsymbol{\mu} \rangle}$ from $\boldsymbol{\mu}$, where $\langle \cdot, \cdot \rangle$ represents the inner product of two vectors. The value s is firstly defined as the mixture rate of two pdfs. However, for a given sample $\mathbf{x}$, we can give another meanings for the value $s(\mathbf{x}) = \frac{\langle 2\mathbf{x} - (\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2) | \mathbf{x} \rangle}{\langle \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1 | \mathbf{x} \rangle}$ as a "label likelihood" of the $\mathbf{x}$, which gives the suitable pdf for the $\mathbf{x}$.

In our new framework, the imbedding mapping and the reduction mapping is set to $\mathbf{x} = \mathbf{f}(s)$ and $s = h(\mathbf{x})$, respectively, as follows:

$$\mathbf{x} = \mathbf{f}(s) = \frac{\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2}{2} + s\frac{\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1}{2} \in S^2, \quad (4)$$

$$s = h(\mathbf{x}) = \frac{\langle 2\boldsymbol{\mu} - (\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2) | \mathbf{x} \rangle}{\langle \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1 | \mathbf{x} \rangle} \in S^1. \quad (5)$$

As mentioned before, there holds $s = h(\mathbf{f}(s))$ when there are no noise. However, the channel is contaminated by noise, and the observed signal $\mathbf{x}$ equals to the sum of the truth $\mathbf{x}_0$ and the noise $\mathbf{m}$ as $\mathbf{x} = \mathbf{x}_0 + \mathbf{m}$. Therefore, an observed (transmitted) label likelihood is the sum of the true label s and the noise n as $s = s_0 + n$. Then, discrete label $s \in \{-1, 1\}$ is transmitted to a continuous label likelihood $s \in S^1 = \mathbb{R}$.

We assume the distribution of $\mathbf{m}$ does not depend on the points in S^2 . This assumption corresponds the homoscedasticity assumption in FLDA (All variance-covariance matrix of classes could be replaced to the same variance-covariance matrix by Eq. (1)). We also assume that noise $\mathbf{m}$ follows a Gaussian distribution with a mean $\mathbf{0}$ and a variance-covariance matrix V . Under these assumptions, the Gramian matrix in S^2 is computed as $F = \frac{1}{2}V^{-1}$.

From Eq. (4)-(5), the tangent matrix T and the reduction matrix R are obtained as

$$T = \frac{\partial \mathbf{f}}{\partial s} = \frac{\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1}{2},$$

$$R = \frac{\partial h}{\partial \mathbf{x}} = \frac{1}{\langle \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1 | \mathbf{x} \rangle} (4\mathbf{x} - \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1)^\top - \frac{\langle 2\mathbf{x} - \boldsymbol{\mu}_1 - \boldsymbol{\mu}_2 | \mathbf{x} \rangle}{\langle \boldsymbol{\mu}_2 - \boldsymbol{\mu}_1 | \mathbf{x} \rangle^2} (\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1)^\top.$$

Here, R depends on s and s should be determined so as to maximize the channel capacity of the label transmission. Consequently, the optimal reduction matrix $\tilde{R}$ is given as

$$\tilde{R} = 2\tilde{G}^{-1}T^\top F = \frac{1}{2\det \tilde{G}} (\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1)^\top V^{-1}.$$

By comparing $(2\det \tilde{G})\tilde{R}^\top = V^{-1}(\boldsymbol{\mu}_2 - \boldsymbol{\mu}_1)$ with Eq. (2), the maximization of channel capacity in transmitting class labels results in FLDA since $V = \Sigma_W$ and $\boldsymbol{\mu}_c = \boldsymbol{\mu}^{(c)}$ ($c = 1, 2$). Then it is geometrically proved that two class FLDA decodes communicating classes to maximize the channel capacity.

5 Conclusion

This paper gives the geometry of two class Fisher's linear discriminant analysis from the point of view of communication theory. As a result, Fisher's linear discriminant analysis gives the maximum channel capacity with transmitting class labels.

References

- [1] S. Amari, Theory of Information Spaces — "A Geometrical Foundation of the Analysis of Communication Systems," RAAG Memoirs, 4:373-418, 1968.
- [2] G. Baudat and F. Anouar, "Generalized Discriminant Analysis Using a Kernel Approach," Neural Computation, 12(10):2385-2404, 2006.
- [3] P. N. Belhumeur et al., "Eigenfaces vs Fisherfaces: Recognition using class specific linear projection," IEEE TPAMI, 19:711-720, 1997.
- [4] T. M. Cover and J. A. Thomas, *Elements of Information Theory (Wiley Series in Telecommunications and Signal Processing)*, Wiley-Interscience, 2006.
- [5] H. P. Decell and S. M. Mayekar, "Feature Combinations and the Divergence Criterion," Computers and Math. with Applications, 3:71-76, 1977.
- [6] R. O. Duda and P. E. Hart, *Pattern Classification and Scene Analysis*, John Wiley and Sons, 1973.
- [7] R. A. Fishier, "The use of multiple measurements in taxonomic problems," Ann. of Eigenics, 7:179-188, 1936.
- [8] J. Fujiki, "Finding discriminant axes from multiple viewpoints," In Proc. the 13th IAPR in Int'l Conf. on Mach. Vision Applications, 73-76, 2013.
- [9] N. Gkalelis, V. Mezaris and I. Kompatsiaris, "Mixture subclass discriminant analysis," IEEE Signal Processing Lett. 18(5):319-322, 2011.
- [10] T. Hastie, A. Buja and R. Tibshirani, "Penalized Discriminant Analysis," The Ann. of Statist., 23(1):73-102, 1995.
- [11] T. Hastie and R. Tibshirani, "Discriminant Analysis by Gaussian Mixture," J. R. Soc. of Statist. Soc. B., 58:155-176, 1996.
- [12] M. Loog, and R. P. W. Duin, "Linear dimensionality reduction via a heteroscedastic extension of LDA: the Chernoff criterion," IEEE TPAMI, 26(6):732-739, 2004.
- [13] H. Sakano et al. "Extended Fisher Criterion Based on Auto-correlation Matrix Information," In Proc. of Structural, Syntactic, and Statist. Patt. Recognit. - Joint IAPR Int'l Workshop, SSPR& SPR 2012, 409-416, 2012.
- [14] A. Sierra, "High-order Fishers discriminant analysis," Patt. Recognit., 35(6):1291-1302, 2002.
- [15] M. Zhu, A. M. Martinez, "Subclass discriminant analysis," IEEE TPAMI, 28(8):1274-1286, 2006.