

최소 가능한 얼굴 인식을 지원하는 치환 변환 기법에 대한 고찰*

김 군 순,^{1*} 강 전 일,¹ 양 대 현,^{1*} 이 경 희²

¹인하대학교 정보보호연구실, ²수원대학교 전자공학과

Revisiting Permutation Transformation Scheme for Cancelable Face Recognition

KoonSoon Kim,^{1*} Jeonil Kang,¹ DaeHun Nyang,^{1*} KyungHee Lee²

¹Information Security Research Laboratory, INHA University,

²Department of Electrical Engineering, The University of Suwon

요 약

생체 정보를 사용하는 시스템에 암호학적 단방향 함수를 직접적으로 적용하는 것은 어려운 것으로 알려져 있다. 이러한 문제를 해결하기 위한 한 가지 방법으로써, 치환 변환 기법이 존재한다. 그러나 그 기법에서는 실험을 통한 구체적인 알고리즘이나 변환 기법에 따른 성능 분석을 직접 보여주지 않았다. 이 논문에서는 변환 기법의 인식률을 실험을 통하여 보여줌으로써 기법이 올바르게 동작한다는 사실을 보인다. 또한 실험 결과를 보여주는데 있어서 LDA에 치환 변환 기법을 적용하였다. 반대로 우리는 또한 치환 변환 기법에 반하는 새로운 공격에 대해서 소개하고, 마지막으로 이 문제를 해결할 수 있는 치환 변환 기법의 일반화를 간략히 소개한다.

ABSTRACT

It is known to be hard to apply cryptographic one-way functions to the recognition system using bio-information directly. As one of the solutions about that problem, there is a permutation transformation scheme. However, they did not show any algorithmic behavior or any performance analysis of the transformation by experiment. In this paper, by showing the recognition ratio of the transformed scheme by experiment, we prove that that scheme is sound. Also, we adopt their transformation to LDA(Linear Discriminant Analysis) to show the experimental results. In the negative side, we introduce a new type of attack against the permutation transformation schemes. Finally, we briefly mention a generalization of the permutation transformation for countermeasure of the attack at the end of this paper.

Keywords : *Biometric, Cancelable, Eigenface, PCA, LDA*

1. 서 론

변환된 도메인의 결과와 원래의 도메인의 결과 비교가 동일하다면, 일방향 함수의 사용은 암호화된 정보의 보안성을 강화시킨다. 그렇기 때문에 많은 인증 기법이 일방향 특성을 사용한다. 생체 인식 분야에서 있어서

* 본 연구는 정보통신부 및 정보통신연구진흥원의 대학 IT 연구센터 지원사업의 연구결과로 수행되었음 (IITA-2006-C1090-0603-0028)

접수일: 2006년 7월 19일 ; 채택일: 2006년 11월 27일

† 주저자, soony@seclab.inha.ac.kr

‡ 교신저자, nyang@inha.ac.kr

일방향성의 특성은 '최소 가능한 특성(cancelable feature)^[14,13]이라 일컫는데, 이는 보안상의 문제가 발생하였을 때 해당 템플릿이 취소되고 새로운 템플릿으로 대체되어야 하기 때문이다. 그러나 생체 인식 기술에서는 변환된 도메인에서의 결과와 원래의 도메인의 비교 결과가 같지 않기 때문에, '최소 가능한 특성'의 필요성이 오랫동안 제기되었음에도 [5]과 [9]와 같은 단지 몇몇의 기법만이 연구 되었다.

IWBRIS 2005에서 발표된 논문 [5]은 특성 얼굴 공간(eigenface space) 위에서 수행되는 취소 가능한 얼굴 인증 기법을 소개하였다. 이 기법은 특성 벡터(eigen vector)와 입력된 사용자 이미지의 치환(permutation)이 두 벡터간의 거리(distance)를 변화 시키지 않는 점을 이용한다. 그러나 논문 [5]에서는 실험을 통한 구체적인 알고리즘이나 변환 기법에 따른 성능 분석을 직접 보여주지 않았다.

이 논문은 다음과 같이 구성된다. II 장에서는 논문 [5]에서 사용한 얼굴 인식 기법인 PCA(Eigenface)와 이 논문에서 추가적으로 사용한 얼굴 인식 기법인 LDA에 대해서 간략히 설명한다. III 장에서는 논문 [5]에서 제시한 취소 가능한 얼굴 인증 기법을 지원하는 치환 변환 기법에 대한 소개를 하고 원리를 설명한다. IV 장에서는 각기 다른 얼굴 인식 기법과 거리 측정 방법을 적용한 치환 변환 기법의 구체적인 실험 결과를 보여준다. 그리고 V 장에서 생체 정보 데이터베이스(biometrics database)로부터 특정 사용자의 템플릿을 추출할 수 있는 새로운 종류의 공격에 대해서 소개하고 이에 대한 대응책을 제시한다.

II. 얼굴 인식 기법의 데이터 분석 기법

1. PCA(Principal Component Analysis): Eigenface

PCA(Principal Component Analysis)와 특성 얼굴(eigenface)을 이용한 얼굴 인식 기법^[11,12]은 1991년에 제안되었다. 이 기법은 PCA가 얼굴 인식에 있어서 좋은 인식률을 보이는 것은 아니지만, 얼굴 이미지를 분류(classify)하기 위하여 주성분(principal components)을 사용한다.

$X = \{x_1, x_2, \dots, x_N\}$ 를 N 개의 샘플 이미지의 집합이라고 하자. 각 샘플 이미지는 m -차원의 벡터로 나타낼 수 있다. $U \in \mathbb{R}^{m \times n}$ 가 n -차원의 이미지 공

간을 m -차원의 이미지 공간으로 변환시키는 행렬일 때 새로운 m -차원의 특징 행렬(feature matrix) $Y \in \mathbb{R}^{m \times N}$ 를 다음과 같이 표현할 수 있다.

$$Y = U^T X \quad (1)$$

PCA에서 $\mu \in \mathbb{R}^n$ 는 모든 샘플 이미지들의 평균이고, 전체 분산 행렬 S_T 를 다음과 같이 정의한다.

$$S_T = \sum_{i=1}^N (x_i - \mu)(x_i - \mu)^T \quad (2)$$

따라서 계산하려는 U 는 다음의 변환된 특징 행렬 $|U^T S_T U|$ 를 최대화하는 행렬이다.

$$U_{opt} = \arg \max_u |U^T S_T U| = [u_1, u_2, \dots, u_m] \quad (3)$$

u_i 는 S_T 의 m 개의 큰 특성값(eigenvalue)과 일치하는, n -차원의 특성 벡터 집합이다.

2. LDA (Linear Discriminant Analysis)

일반적으로 LDA (Linear Discriminant Analysis)는 다른 클래스간의 분산 정보(between-class)와 같은 클래스의 데이터에 존재하는 분산 정보(within-class)을 이용하기 때문에, PCA보다 데이터들의 분류(classification)에 더욱 적합하지만, SSS(Small Sample Size) 문제가 있다.

SSS(Small Sample Size) 문제는 샘플 이미지의 차원보다 샘플 이미지의 갯수가 적은 경우 발생하며, 이를 해결하기 위하여 Fisherface^[15], F-LDA^[17], D-LDA^{[8][4][6]}, R-LDA^[7]와 같은 다수의 방법이 존재한다.

$X = \{X_i\}_{i=1}^C$ 이 샘플 클래스들의 집합이라고 하고, $X = \{X_{ij}\}_{j=1}^{C_i}$ 이 C_i 샘플 이미지들의 클래스라고 하면, 샘플 이미지들의 전체 수는 $N = \sum_{i=1}^C C_i$ 이다. 각각의 샘플 이미지는 n -차원의 벡터로 나타낼 수 있고, PCA에서처럼, n -차원의 이미지 공간을 m -차원의 이미지 공간으로 변환하기 위한 행렬 $U \in \mathbb{R}^{n \times m}$ 를 정의할 수 있다.

LDA에서 μ 는 전체 샘플 공간 X 의 평균을 나타내고, μ_i 는 i -번째 클래스 X_i 의 평균을 나타낸다. 클

래스들 간의 분산 행렬(between-class scatter matrix) S_b 와 클래스 내의 분산 행렬(within-class scatter matrix) S_w 는 다음과 같이 정의한다.

$$S_b = \frac{1}{N} \sum_{i=1}^C C_i (\mu_i - \mu) (\mu_i - \mu)^T \quad (4)$$

$$S_w = \frac{1}{N} \sum_{i=1}^C \sum_{j=1}^{C_i} (x_{ij} - \mu_i) (x_{ij} - \mu_i)^T \quad (5)$$

U 는 클래스 내의 분산 행렬 분의 클래스들 간의 분산 행렬의 비율을 최대로 만듦으로써 구할 수 있다.

$$U_{opt} = \arg \max_U \frac{|U^T S_b U|}{|U^T S_w U|} = [u_1, u_2, \dots, u_m] \quad (6)$$

R-LDA에서, U 는 $0 \leq \eta \leq 1$ 는 클래스들 간의 분산 행렬에 대한 가중치이고, 다음과 같이 계산할 수 있다.

$$U_{opt} = \arg \max_U \frac{|U^T S_b U|}{|\eta (U^T S_b U) + U^T S_w U|} \quad (7)$$

III. 치환 변환 기법

1. 치환 변환 기법^{[1][5]}의 주요 개념

μ 는 전체 샘플 이미지들의 평균이고, x 는 테스트 이미지이다. 각 샘플 이미지 x_i 는 아래의 어떠한 행렬 U 의 $y_i \in \mathbb{R}^m$ 로 나타낼 수 있다.

$$y_i = U^T (x_i - \mu) \quad (8)$$

이와 유사하게 테스트 이미지는 대표 $y \in \mathbb{R}^m$ 로 나타낼 수 있다.

$$y = U^T (x - \mu) \quad (9)$$

y_i 와 y 간의 모든 거리를 계산한 후, y_i 와 y 간의 거리가 가장 근접한 값을 찾아서 테스트 이미지 x 에 대한 j -번째 이미지를 구할 수 있다. 이 때 만약 어

떠한 한계치 값 θ_ϵ 에 대해서 $\|y_i - y\| > \theta_\epsilon$ 라면, 테스트 이미지 x 를 '알 수 없음'으로 분류한다.

논문 [1][5]에서 언급된 치환 변환 기법은 행렬 연산의 성질을 이용한다. $f_c()$ 와 $f_r()$ 은 어떠한 입력 행렬을 동일한 방법으로 각각 열(column)과 행(row)으로 치환하는 함수라고 할 때, 다음과 같은 식을 만족한다.

$$AB = f_c(A) f_r(B) = (A P_c) \times (P_r B) \quad (10)$$

$f_c()$ 와 $f_r()$ 은 어떠한 치환 행렬 P_c 와 P_r 을 입력 행렬의 뒤와 앞에 각각 곱하게 되는데, 이 때 곱해지는 치환 행렬들은 PBKDP(Password-Based Key Derivation Function)^[16]를 이용하여 '사용자 키워드(user keyword)*'로부터 생성할 수 있다. 따라서 얼굴 이미지 x 와 사용자 키워드가 일치한다면 다음의 식을 만족하게 된다.

$$y = U^T (x - \mu) = f_c(U^T) f_r(x - \mu) = y' \quad (11)$$

일치하지 않는 경우에는 다음의 식을 만족한다.

$$y = U^T (x - \mu) \neq f_c(U^T) f_r(x - \mu) = y' \quad (12)$$

각 사용자에게 대한 $f_c(U^T)$ 와 $f_r(\mu)$ 의 연산 결과는 템플릿의 형태로 시스템에 저장된다. 이 데이터들은 사용자 키워드를 바꿈으로써 시스템의 데이터베이스 내에서 취소가 가능하다. 물론 이 때, 어떠한 저장 공간 절약 메커니즘(storage-saving mechanism)이 사용되었다.

논문 [1][5]에서는, PCA와 특성 얼굴 공간(eigenface space)을 이용하여 작동하는 기법으로 제안되었지만, 위의 치환 연산이 어떤 행렬 U 로 대표되는 샘플 이미지에 적용되기 때문에, 행렬 U 를 만드는 방법에 있어서 반드시 PCA 방식이 사용될 필요는 없다. 따라서 얼굴이나 다른 무언가의 인식을 위해서 분류(classification)를 수행할 목적으로 의미 있는 특성 벡터를 얻어낼 수 있다면, 치환 변환 기법은 다른 데이터 분류 기법과도 연동이 가능하다.

이 기법은 안전한 치환 행렬을 생성하기 위해서, 사용자가 시스템에 의해 인증 받기를 원할 때마다

* 논문 [1][5]에서 사용자 패스워드로 언급되었으나, 이 논문에서는 패스워드 대신에 키워드라는 명칭을 사용한다. 패스워드는 보안성의 의미를 강하게 내포하기 때문이다.

입력하는 사용자 키워드를 사용한다. 이것은 사용자 측면에서는 다소 불편할 수 있겠지만, 취소 가능한 특성(cancelable feature)을 지원하는 모든 기법들은 어떠한 형태로든 생체정보 이외의 정보를 필요로 한다. 생체정보 갱신을 위한 이러한 정보를 인증 시에 사용자로부터 입력 받거나, 시스템에 저장하고 있을 수 있다. 예를 들어서 Fuzzy Vault 기법을 이용한 생체 인증 기법은 인증 시에 사용자로부터 직접적으로 어떠한 참여정보를 필요로 하지 않지만, Vault를 생성하기 위한 특정 정보를 시스템에 저장할 경우 생체정보가 공격자에게 노출되면 인증에 대한 위조가 가능하다. 이에 반해 치환 변환 기법에서는 만약 각 사용자의 키워드가 비밀로써 사용되어 공격자에게 노출되지 않는다면, 공격자는 성공적인 인증을 위한 치환 행렬을 생성하기 위해서 이 비밀이 필요하기 때문에, 공격자는 생체 정보만을 알고 있어서는 인증을 위조할 수가 없다.

어떠한 치환을 수행할 것이냐는 관리자가 결정할 문제지만, 모든 사용자에 대해서 동일한 치환을 사용하는 것은 각 사용자의 보안 측면에는 좋지 않다. 왜냐하면 공격자가 하나의 치환을 알아낸다면, 공격자는 모든 사용자의 템플릿을 얻어낼 수 있기 때문이다. 취소 가능한 생체 보안 기술 중 하나인 논문 [9]에서, MACE 필터(filter)를 생성하기 위해 PIN(personal identification number)을 사용한다. 또한 바이오 해시(biometric hash) 논문 [3]에서는, 사용자 고유의 토큰(token)을 사용하여 동일한 사용자의 얼굴 이미지로부터 인증을 위한 고유한 키를 얻어내며, 토큰의 변경을 통하여 취소 가능한 특성을 지원한다. 치환 변환 기법은 사용자 고유의 임의 행렬을 고유의 얼굴 벡터에 적용시키지만, 바이오 해시 기법은 사용자 고유의 토큰을 사상된 벡터(projected vector)에 적용시킨다. 바이오 해시 기법의 장점은 시스템에 어떠한 생체 정보도 남기지 않는다는 것이다. 이에 반해, 치환 변환 기법은 인증 과정을 위하여 치환된 생체 정보를 시스템에 유지해야 하므로 보안상의 문제를 가질 수 있다. 그러나 바이오 해시 기법은 치환 변환 기법에 비하여 두 세배 큰 사상 행렬(projection matrix)을 유지해야 한다.

2. 인식률에 대한 논의

치환 변환 기법은 $P_c = P_r^T$ 와 $P_c P_r = I$ 의 성질을

지니는 두 개의 치환 행렬 $P_c, P_r \in \mathbb{R}^{n \times n}$ 을 사용한다. 또한 모든 정수 $1 \leq j \leq n$ 와 $x \in \{c, r\}$ 에 대하여 $\sum_{i=1}^n P_x[i][j] = 1$ 와 $\sum_{i=1}^n P_x[j][i] = 1$ 인 성질을 만족한다.

잘못된 사용자 키워드로부터 생성된 치환 P_r' 을 가정해 보자. 사용자가 잘못된 키워드를 입력하면 다음의 식을 만족한다.

$$\begin{aligned} y' &= U^T P_c (P_r' x - P_r \mu) \\ &= U^T P_c P_r' x - U^T P_c P_r \mu \\ &= U^T P_c P_r' x - U^T \mu \end{aligned} \quad (13)$$

여기서 $P_c P_r'$ 이 다른 치환 행렬이 됨을 쉽게 알 수 있다. 정상적인 경우와 비교할 때, y' 는 다음과 같은 추가적인 노이즈를 포함하고 있다.

$$\begin{aligned} y' - y &= (U^T P_c P_r' x_i - U^T P_c P_r \mu) - (U^T x_j - U^T \mu) \\ &= U^T P_c P_r' x_i - U^T x_j \\ &= U^T (P_c P_r' x_i - x_j) \end{aligned} \quad (14)$$

이 때 만약 i 와 j 가 같다면, 노이즈는 오로지 사용자 키워드의 차이에 따라서 결정된다.

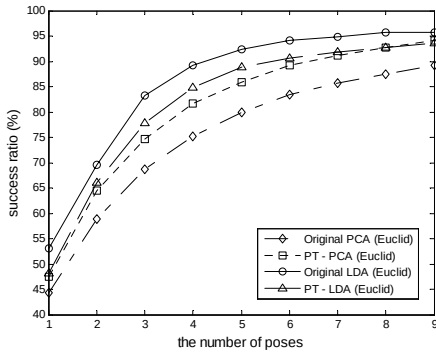
$$P_c P_r' x_i - x_j = P_c (P_r' - P_r) x_i \quad \text{if } i = j \quad (15)$$

논문 [1][5]에서는, 노이즈가 발생하는 경우는 고려하지 않음으로써 동일한 인식률을 보일 것으로 예상했지만, 이 논문에서는 노이즈의 발생이 얼굴 인식률에 영향을 줄 수 있을 것으로 생각한다.

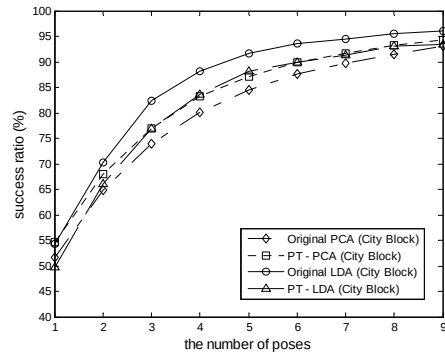
IV. 실험

1. 실험 조건

R-LDA 기법에서 n 의 값은 샘플 집합의 자세의 수(the number of poses)가 2보다 작거나 같은 경우에 1로 설정하고, 2보다 큰 경우에는 0.001(≈ 0)로 설정하여 LDA 실험을 수행 하였다. 실험에 사용한 얼굴 데이터베이스는 ORL(Olivetti Research Lab) DB이다. ORL DB는 40명의 각기 다른 사람으로 나뉘고, 각 사람별로 10종류의 다른 자세(pose)의 이미지를 포함하고 있다. 그리

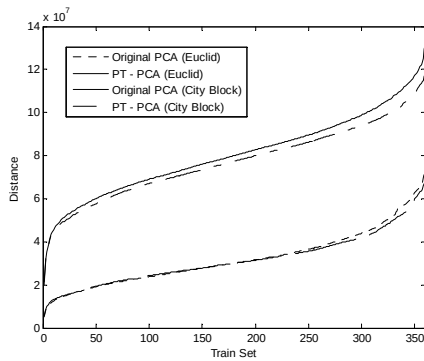


(a) 유클리드 거리(Euclid Distance)

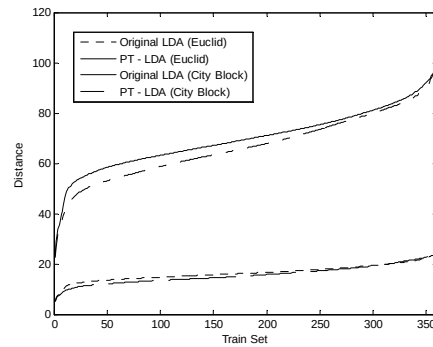


(b) 시티 블록 거리(City Block Distance)

(그림 1) PCA, LDA, PT-PCA 그리고 PT-LDA의 각 기법에 대한 얼굴 인식률의 비교



(a) PCA 경우



(b) LDA 경우

(그림 2) (a) PCA와 (b) LDA에 치환 변환 기법이 적용된 경우에, 하나의 샘플 이미지에서 다른 360개 샘플 이미지들 간의 정렬된 거리

고 각 샘플 이미지는 64×64 픽셀(pixel)의 크기로 얼굴 부분만을 편집한 이미지를 사용하였다. PCA 실험에서는 U^T 에 대한 특성 벡터의 수 m 은 40으로 설정하였다.

이 논문에서 사용한 벡터간 거리 측정 방법은 유클리드 거리(Euclid Distance)와 시티 블록 거리(City Block Distance)이다. 이 때 $a, b \in \mathbb{R}^k$ 이다. 유클리드 거리 방법은 다음과 같이 정의한다.

$$d(a, b) = \|a - b\|^2 = \sum_{i=1}^k (a_i - b_i)^2 \quad (16)$$

그리고 시티 블록 거리는 다음과 같이 정의한다.

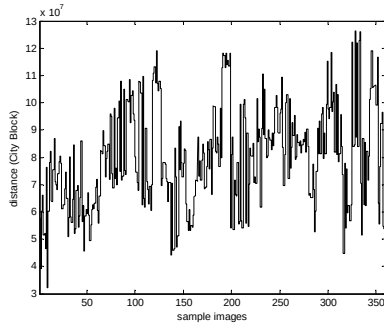
$$d(a, b) = |a - b| = \sum_{i=1}^k |a_i - b_i| \quad (17)$$

2. 실험 결과

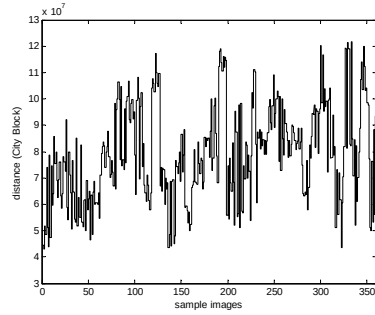
[그림 1]에서 x 축은 학습 특성 벡터(training eigenvector) U^T 의 자세들의 수를 의미하고 y 축은 테스트 이미지들을 인식한 성공률을 의미한다. 각 자세에서의 테스트는 100번 이상 수행되었다. 또한 각 테스트 단계마다, DB에서 학습 이미지들을 제외한 모든 이미지들이 테스트 집합으로서 실험에 사용되었다. 보는 바와 같이, 유클리드 거리 대신에 시티 블록 거리를 사용하여 실험한 경우 인식률이 3~5%정도 높게 나타난다.

PCA 기법에 적용한 치환 기법 실험에서는 다소 높은 인식률을 보여주었으나, LDA 기법에서의 실험에서는 다소 저조한 인식률을 나타냈다.

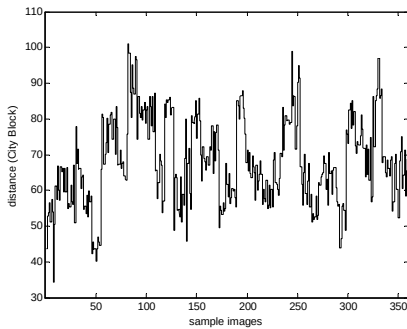
이러한 결과의 이유는 치환 기법이 샘플 벡터의



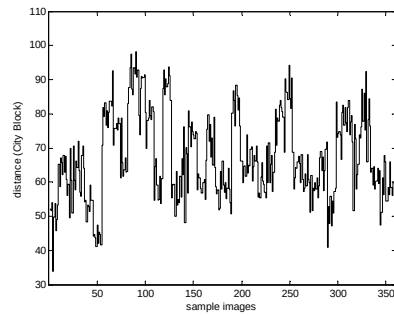
(a) PCA에 치환 변환 기법이 적용된 경우 1



(b) PCA에 치환 변환 기법이 적용된 경우 2



(c) LDA에 치환 변환 기법이 적용된 경우 1



(d) LDA에 치환 변환 기법이 적용된 경우 2

[그림 3] 치환 변환 기법이 (a)(b) PCA와 (c)(d) LDA에 적용된 경우에, 하나의 샘플 이미지에서 다른 360개의 샘플 이미지들 간의 거리. 각각의 경우에서, 각기 다른 경우의 사용자 키워드들을 모든 샘플 이미지들에 적용하였다. (벡터간의 거리 측정 방법은 시티 블록 거리를 사용함)

방향을 바꾸는 것과 관련되어 있다. 이러한 특성은 다음 절에서 언급될 새로운 공격 방법과 관련 있다. 바꾸어 말하면, PCA에서 치환 기법의 인식률이 좋게 나오는 이유는 기존의 PCA에서의 샘플 벡터 방향이, 치환에 의해 임의적으로(randomly) 분배되었을 때보다 나쁘기 때문이다. 이는 논문 [10]의 감독되지 않은 학습 방법(unsupervised learning method)이 PCA 기법에서 인식률이 더 높게 나타나는 것과 유사하다. 그러나 LDA에서 치환 변환의 임의성(randomness)은 이미 최적화 된 상태를 망치게 되어 인식률을 떨어뜨리는 결과를 초래한다. PCA와 LDA의 두 기법의 경우에 치환 기법이 적용될 때, 임의의 한 개의 샘플 이미지에서 다른 샘플 이미지까지의 거리는 전체적으로 줄어든다.

[그림 2]에 제시된 바에 의하면, PCA 기법에서 나타난 결과는 위에서 언급된 인식률을 따르지 않는 것처럼 보이지만, 거리가 짧은 경우에서의 거리 값들이 기존의 PCA 기법보다 높다. 이것은 치환 기법

이 PCA에서 유사한 샘플 이미지들과 테스트 이미지 사이의 거리를 더 높게 만든다는 것을 의미한다. 벡터의 방향이라는 측면에서, 치환 기법은 유사한 샘플 이미지를 테스트 이미지와는 다른 방향으로 만든다는 의미를 지닌다.

치환 변환 기법에서 발생하는 노이즈에 대한 가정은 [그림 3]에 의해 간접적으로 증명된다. 우리는 모든 이미지들에 대하여, 임의적으로 선택된 다른 사용자 키워드들을 적용하여 실험에 진행하였으나, 도출된 그래프의 모양은 너무나 유사했다. 이것은 거리를 구하는데 있어서, 벡터의 크기가 벡터의 방향성보다 더욱 영향을 미친다는 것을 의미한다. LDA 기법은 실험 결과의 그래프에서 볼 수 있듯이, 벡터의 방향에 의해 더 많은 영향을 받는다.

V. 특정 사용자의 템플릿을 추출하는 공격 방법

논문 [1][5]에서, 제시한 기법에는 다양한 관점

에서 그들의 보안성을 기술하였다. 대부분은 매우 설득적이지만, 이 장에서는 템플릿 자체의 거리를 이용하여 특정 사용자 템플릿을 추출하는 공격에 대해 생각해 볼 수 있다.

1. 템플릿의 거리

보다 높은 인식률을 얻기 위해서, 두 개 벡터들의 차 벡터(difference vector)가 사용되지만, 공격자가 반드시 이 방식을 따라야할 필요는 없다. 그림 4에서 보는 바와 같이, 만약에 거리 값들의 차를 식별하기 쉬운 사용자가 있다면, 공격자는 공격을 위해 오직 순수한 거리 값만을 사용할 수 있기 때문이다.

치환 변환 기법에서, 치환 행렬 P_c 와 P_r 를 사용하였다. 공격자가 데이터베이스를 공격하여, 저장된 템플릿 $\langle U_{RR}^T, \mu_{RR}, S_{x,i} \rangle$ 와 저장된 대표 이미지 y_i 를 얻을 수 있다면, 다음의 식에 나타난 것처럼 공격자는 치환된 이미지 $P_{r,i}x_i$ 를 추출할 수 있을 것이다.

$$y_i = U_{RR}^T P_{c,i} \{P_{r,i}x_i - P_{r,i}\mu\} \quad (18)$$

$$= U_{RR}^T S_{c,i} \{P_{r,i}x_i - S_{r,i}\mu_{RR}\}$$

$$P_{r,i}x_i = (U_{RR}^T S_{c,i})^{-1} y_i + S_{r,i}\mu_{RR} \quad (19)$$

공격자가 특정 사용자의 얼굴 이미지 x_i 를 가지고 있다면, $|P_{r,i}x_i| \simeq |x_i|$ 를 점검함으로서 데이터베이스

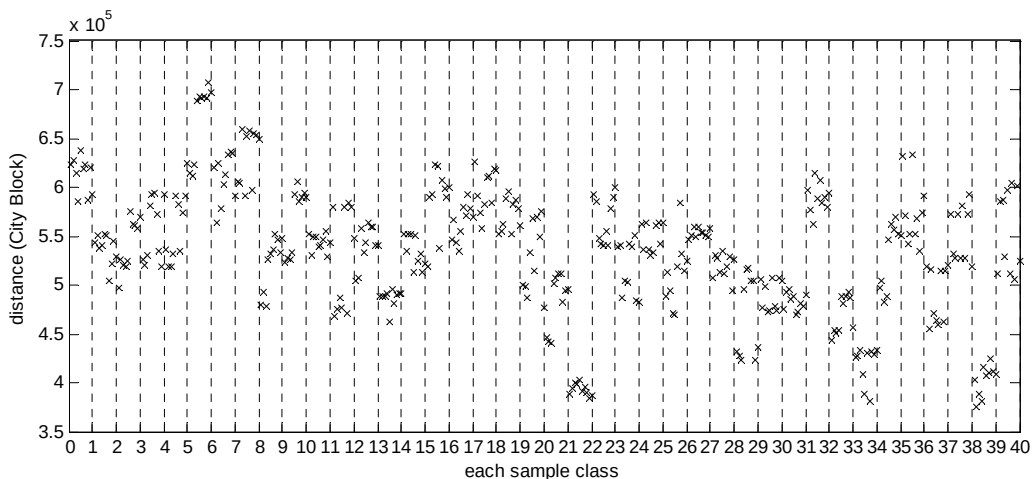
스 내에 일치된 사용자가 등록되어있는지 여부를 알 수 있다. 물론 U^T 가 샘플 이미지들의 특성 벡터의 전체를 포함하지 않는다는 사실은 이러한 공격에 도움이 될 수 있지만, 큰 영향을 미치지는 않는다.

일반적으로 얼굴 이미지들의 거리는 얼굴의 특성으로부터 직접적인 영향을 받으므로, 공격 대상인 사람의 얼굴이 특이한 경우에, 공격자는 보다 쉽게 공격할 수 있을 것이다.

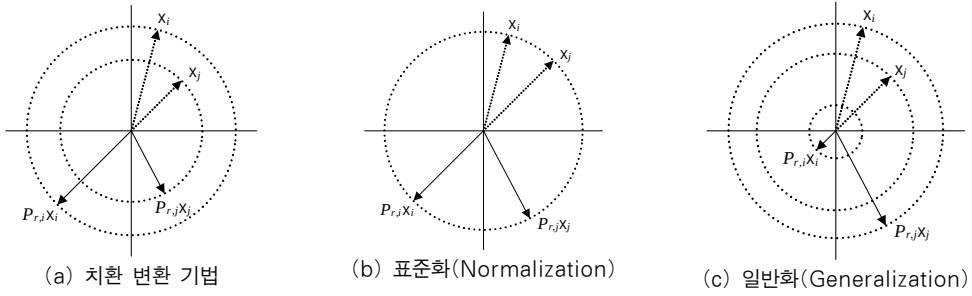
이러한 공격이 가능한 이유는, 치환 행렬 P_c 와 P_r 이 곱해진 행렬의 거리를 바꾸기 않고, 단지 행렬의 방향만 바꾸기 때문이다. 이에 관한 논의는 지난 장에 나타난 인식률에 관련되어 있다. 즉, 원래의 방향과 어떠한 거리와 유사하다면 인식률은 증가하나, 원래의 방향이 다른 것과 현저히 다른 경우 인식률은 떨어진다는 것을 의미한다. 이것이 LDA 기법에서 채택한 치환 변환 기법이 기존의 LDA 기법보다 다소 낮은 인식률을 보이는 이유이며, LDA 기법이 PCA 기법보다 구별에 보다 적합한 이유이다.

2. 공격에 대비한 대응책

이 논문에서는, 위에서 언급한 취약점을 해결하기 위해 두 가지의 대응책을 제시한다. 하나는 표준화(normalization)이고, 다른 하나는 일반화(generalization)이다. 표준화는 모든 샘플 이미지들의 거리를 동일하게 만드는 방법이다. 따라서 거리를 이용한 공격이 무의미해진다. 표준화의 기본 개



[그림 4] 원점에서 각 샘플 이미지들간의 거리. 각 클래스의 샘플 이미지는 동일한 모노(grid)내에 위치한다.



[그림 5] 2차원 공간상에 그려진 표준화와 일반화의 개념도

넘은 [그림 5]에서 볼 수 있다.

그러나 샘플 이미지들의 거리가 주로 샘플의 인식에 사용된다는 점에서, 표준화는 인식률에 부정적인 영향을 미칠 가능성이 있다. 달리 말하면, 표준화의 도입은 인식률을 떨어뜨릴 수 있으며, 더욱이 외모가 특이한 사람의 경우 보안성 대신에 더 낮은 인식률을 보일 것이다.

두 번째 방법인 일반화는 조금 더 복잡하다. 기존의 치환 변환 기법은 $|U^T P_c| = |U^T|$ 와 $|P_r(x_i - \mu)| = |x_i - \mu|$ 의 성질을 만족한다. 만약 잘못된 사용자 키워드로부터 생성된 P_r' 이 치환 변환에 적용되었을 시에, $|U^T P_c| \neq |U^T|$ 와 $|P_r(x_i - \mu)| \neq |x_i - \mu|$ 의 성질을 만족하도록 P_c 와 P_r 을 변경할 수 있다면, $|P_r' - P_r|$ 이 작은 값을 갖지 않을 것이다. 즉, 추가되는 노이즈가 더 커지게 된다. 이를 위하여, $\langle P_c, P_r \rangle$ 을 생성하는 방법을 다시 설계할 필요가 있다.

치환 변환 기법에서 사용되는 치환 행렬 P_c 와 P_r 는 $P_c P_r = I$ 을 만족하는 일반적인 행렬로 변환될 수 있다. 그렇다면 하더라도 $AB = (AP_c) \times (P_r B)$ 의 성질을 여전히 만족하게 되며, 따라서 $|P_{r,i}x_i|$ 와 $|x_i|$ 를 비교하는 것은 더 이상의 의미를 갖지 않게 된다.

노이즈가 충분히 커진다면, 인식률은 증가하게 될 것이다. 그러나 이 방법에서, P_c 와 P_r 의 곱셈 연산은 많은 메모리 공간과 더 많은 시간을 필요로 한다. 왜냐하면 두 행렬간의 곱셈 연산은 실제로 4096×4096 행렬 상에서 수행되기 때문이다. (치환 기법에서 곱셈 연산은 연산의 대상이 되는 행렬들을 직접 치환함으로써 수행된다.)

VI. 결 론

이 논문에서, [1][5]에 언급된 기법은 보안성에서 뿐만 아니라, 인식률에 있어서도 가치가 있다는 것을 실험을 통하여 보여주었다. 또한 추가적으로 LDA 기법을 치환 변환 기법에 적용하였다. 적용한 치환 변환 기법에서의 인식률은 LDA보다는 낮았으나, 본래의 PCA보다는 높았으며, 이에 대한 이유는 자세하게 분석하였다.

또한 논문에서는 원점에서 이미지까지의 거리를 이용한 공격에는 취약한 문제점이 제시하였고, 이를 해결하기 위하여 두 가지 대응책을 생각해보았다. 표준화는 모든 이미지들의 거리를 동일하게 만드는 방법이며, 일반화는 거리를 공격자에게 노출되지 않도록 하는 방법이다.

특히 두 번째 해결방안인 일반화는 주목할 만하다. 이 기법이 예상대로 수행된다면, 일반화의 방법은 LDA에 적용되더라도 인식률이 증가할 것이기 때문이다. 이를 증명하기 위해서는 더 깊은 연구와 다양한 실험이 필요하다.

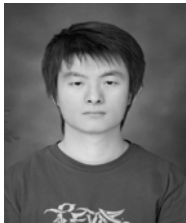
참 고 문 헌

- [1] 강전일, 양대현, 이경희, “두 가지 보안 요소를 사용하는 최소 가능한 얼굴 인증 기술”, 한국정보보호학회 논문지 제 16권 제 1호 pp. 13-21 2006년 2월
- [2] Aleix M. Martinez, and Avinash C. Kak, “PCA versus LDA”, IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol.23,

- No.2, pp. 228-233, Feb. 2001
- [3] David C. L. Ngo, Andrew B. J. Tech, and Alwyn Goh, Biometric Hash: High-Confidence Face Recognition, IEEE Transaction on Circuits and Systems for Video Technology, Vol.16, No.6, June 2006
- [4] Hua Yu and Jie Yang. "A Direct LDA Algorithm for High-Dimensional Data - with Application to Face Recognition", Pattern Recognition, Vol.34, No.10, pp. 2067-2070, Oct. 2001
- [5] Jeonil Kang, DaeHun Nyang and KyungHee Lee, "Two Factor Face Authentication Scheme with Cancelable Feature", the Proceeding of IWBRIS 2005, LNCS 3781, pp. 67-75, Oct. 2005
- [6] Juwei Lu, Kostantinos N. Plataniotis and Anastasios N. Venetianopoulos. "Face Recognition Using LDA-Based Algorithms", IEEE Transactions on Neural Networks, Vol. 14, No.1, pp.195-200, Jan. 2003
- [7] Juwei Lu, Kostantinos N. Plataniotis and Anastasios N. Venetianopoulos. "Regularization Studies of Linear Discriminant Analysis in Small Sample Size Scenarios with Application to Face Recognition", Elsevier Science Inc., Pattern Recognition Letters, Vol.26, Issue 2, pp. 181-191, Jan. 2005
- [8] L. Chen, H. Liao, M. Ko, J. Lin and G. Yu. "A new LDAbased face recognition system which can solve the small sample size problem", Pattern Recognition, Vol.33, No. 10, pp. 1713-1726, 2000.
- [9] Mario Savvides, B.V.K. Vijaya Kumar and P.K. Khosla. "Cancelable biometric filters for face recognition", Pattern Recognition, 2004. ICPR 2004. Proceedings of the 17th International Conference on Vol.3, pp. 922- 925, Aug. 2004
- [10] Marian Stewart Bartlett, Javier R. Movellan and Terrence J. Sejnowski. "Face Recognition by Independent Component Analysis", IEEE Transactions on Neural Networks, Vol.13, No.6, pp. 1450-1464, Nov. 2002
- [11] Matthew A. Turk and Alex P. Pentland. "Face Recognition Using Eigenfaces", Computer Vision and Pattern Recognition, 1991. Proceedings CVPR '91., IEEE Computer Society Conference on 3-6 pp. 586 - 591, June 1991
- [12] Matthew A. Turk and Alex P. Pentland. "Eigenfaces for Recognition", Journal of Cognitive Neuroscience, Vol. 3, No. 1, pp. 71-86, 1991.
- [13] Michael Braithwaite, Ulf Cahn von Seelen, James Cambier, John Daugman, Randy Glass, Russ moore and Ian Scott. "Application-Specific Biometric Templates", Iridian Technologies Inc., Proceedings of AutoID, pp. 167 - 171, 2002
- [14] N.K. Ratha, J.K. Connell and R.M. Bolle, "Enhancing security and privacy in biometrics-based authentication systems", IBM Systems Journal, Vol.40, No.3, pp. 614-634, 2001
- [15] Peter N. Belhumeur, Joao P. Hespanha and David J. Kriegman. "Eigenfaces vs. Fisherfaces: Recognition Using Class Specific Linear Projection". IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol.19, No. 7, pp. 711-720, July 1997

- [16] PKCS#5 v2.0, Password-Based Cryptography Standard, RSA Laboratory, March 25, 1999
- [17] R. Lotlikar and R. Kothari. "Fractional-step dimensionality reduction", IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol.22, pp. 623-627, June 2000
- [18] Wendy S. Yambor, Bruce A. Draper and J. Ross Beveridge. "Analyzing PCA-based Face Recognition Algorithms: Eigenvector Selection and Distance Measures", 2nd Workshop on Empirical Evaluation in Computer Vision, 2000.

〈著者紹介〉



김 군 순 (KoonSoon Kim) 학생회원

2006년 8월 : 인하대학교 컴퓨터 공학과 졸업
 2006년 9월~현재 : 인하대학교 정보통신대학원 석사
 <관심분야> 생체 인식 보안



강 전 일 (Jeonil Kang) 학생회원

2003년 2월 : 인하대학교 컴퓨터 공학과 졸업
 2006년 2월 : 인하대학교 정보통신대학원 석사
 2006년 3월~현재 : 인하대학교 정보통신공학과 박사 과정
 <관심분야> RFID 보안, 생체 인식 보안, 무선 센서 네트워크, 무선 인터넷 보안



양 대 헌 (DaeHun Nyang) 정회원

1994년 2월 : 한국과학기술원 과학기술 대학 전기 및 전자 공학과 졸업
 1996년 2월 : 연세대학교 컴퓨터 과학과 석사
 2000년 8월 : 연세대학교 컴퓨터 과학과 박사
 2000년 9월~2003년 2월 : 한국전자통신연구원 정보보호연구본부 선임연구원
 2003년 2월~현재 : 인하대학교 정보통신대학원 조교수
 <관심분야> 암호이론, 암호프로토콜, 인증프로토콜, 무선 인터넷 보안



이 경 희 (KyungHee Lee) 정회원

1989년 : 서울대학교 식품영양학과 학사
 1993년 : 연세대학교 전산학과 학사
 1998년 : 연세대학교 컴퓨터과학과 석사
 2004년 : 연세대학교 컴퓨터과학과 박사
 1993년 1월~1996년 5월 : LG소프트(주) 연구원
 2000년 12월~2005년 2월 : 한국전자통신연구원 선임연구원
 2005년 3월~현재 : 수원대학교 전임강사
 <관심분야> 영상처리, 컴퓨터비전, 인공지능, 패턴인식, 생체인식, 얼굴인식, 다중생체인식