

Computation of Cryptographic Keys from Face Biometrics

Alwyn Goh¹ and David C.L. Ngo²

¹ Corentix Laboratories, B-19-02 Cameron Towers, Jln 5/58B,
46000 Petaling Jaya, Malaysia.
alwyn_goh@yahoo.co.uk

² Faculty of Information Science & Technology, Multimedia University,
75450 Melaka, Malaysia

Abstract. We outline cryptographic key-computation from biometric data based on error-tolerant transformation of continuous-valued face eigenprojections to *zero-error* bitstrings suitable for cryptographic applicability. *Bio-hashing* is based on iterated inner-products between pseudorandom and user-specific eigenprojections, each of which extracts a single-bit from the face data. This discretisation is highly tolerant of data capture offsets, with same-user face data resulting in highly correlated bitstrings. The resultant user identification in terms of a small bitstring-set is then securely reduced to a *single* cryptographic key via Shamir secret-sharing. Generation of the pseudorandom eigenprojection sequence can be securely parameterised via incorporation of physical tokens. Tokenised bio-hashing is rigorously protective of the face data, with security comparable to cryptographic hashing of token and knowledge key-factors. Our methodology has several major advantages over conventional biometric analysis ie elimination of false accepts (FA) without unacceptable compromise in terms of more probable false rejects (FR), straightforward key-management, and cryptographically rigorous commitment of biometric data in conjunction with verification thereof.

1 Introduction

Biometric ergonomics and cryptographic security are highly complementary attributes, hence the motivation for the presented research. Computation of cryptographic keys from biometric data was first proposed in the Bodo patent [1], and is technically challenging from both signal processing and information security viewpoints. The representation problem is that biometric data (ie linear time-series or planar bitmaps) is continuous and high-uncertainty, while cryptographic parameters are discrete and zero-uncertainty. Biometric consistency—ie the difference between reference and test data, which are (at best) similar but never equal—is hence inadequate for cryptographic purposes which require exact reproduction. This motivates the formulation of offset-tolerant discretisation methodologies, the end result of which is also required to be protect against adversarial recovery of user-specific biometrics.

2 Review of Previous Work

The earliest publications in this domain are by Soutar et al [2, 3], whose research outlines cryptographic key-recovery from the integral correlation of freshly captured fingerprint data and previously registered *bioscrypts*. Bioscrypts result from the mixing of random and user-specific data—thereby preventing recovery of the original fingerprint data—with data capture uncertainties addressed via multiply-redundant majority-result table lookups. This ensures representation tolerance against offsets in same-user test fingerprints, but does not satisfactorily handle the issue of discrimination against different-user data.

The Davida et al [4, 5] formulation outlines cryptographic signature verification of iris data without stored references. This is accomplished via open token-based storage of user-specific Hamming codes necessary to rectify offsets in the test data, thereby allowing verification of the corrected biometrics. Such self-correcting biometric representations are applicable towards key-computation, with recovery of iris data prevented by complexity theory. Resolution of biometric uncertainty via Hamming error correction is rigorous from the security viewpoint, and improves on the somewhat heuristic Soutar et al lookups.

Monrose et al key-computation from user-specific keystroke [6] and voice [7] data is based on the deterministic concatenation of single-bit outputs based on logical characterisations of the biometric data, in particular whether user-specific features are below (0) or above (1) some population-generic threshold. These feature-derived bitstrings are used in conjunction with randomised lookup tables formulated via Shamir [8] secret-sharing. Error correction in this case is also rigorous, with Shamir polynomial thresholding and Hamming error correction considered to be equivalent mechanisms [5]. The inherent scalability of the bitstrings is another major advantage over the Soutar et al methodology.

Direct mixing of random and biometric data (as in Soutar et al) allows incorporation of serialised physical tokens, thereby resulting in token+biometric cryptographic keys. There are also advantages from the operations security viewpoint, arising from the permanent association of biometrics with their owners. Tokenised randomisation protects against biometric fabrication—as demonstrated by Matsumoto et al [9] for fingerprints, which is considered one of the more secure form factors—without adversarial knowledge of the randomisation, or equivalently possession of the corresponding token.

3 Bio-Hash Methodology

This paper outlines cryptographic key-computation from face bitmaps, or specifically from Sirovich-Kirby [10, 11] eigenprojections thereof. The proposed bio-hashing is based on: (1) *biometric eigenanalysis*: resulting in user-specific eigenprojections with a moderate degree of offset tolerance, (2) *biometric discretisation*: via iterated inner-product mixing of tokenised and biometric data, with enhanced offset tolerance, and (3) *cryptographic interpolation*: of Shamir secret-shares corresponding to token and biometric data, culminating in a zero-error key. Bio-hashing has the following ad-

vantages: (1) *tokenised random mixing*: in common with Soutar et al, (2) *discretisation scalability*: in common with Monrose et al, and (3) *rigorous error correction*: in common with Davida et al and Monrose et al. The proposed formulation is furthermore highly generic arising from the proposed discretisation in terms of inner-products ie $s = \mathbf{a} \cdot \mathbf{b}$ for $\mathbf{a}, \mathbf{b} \in \mathbb{R}^n$

We believe our work to be the first demonstration of key-computation from face data, which seems difficult to handle (in common with other planar representations) using the Monrose et al procedure. Bio-hashing is essentially a transformation from representations which are high-dimension and high-uncertainty (the face bitmaps) to those which are low-dimension and zero-uncertainty (the derived keys). The successive representations are: (1) *raw bitmap*: $\mathbf{x} \in S$ in domain $\mathbb{R}^N$, with N the pixelisation dimension, (2) *eigenprojection*: $\mathbf{a} \in S'$ in domain $\mathbb{R}^n$, with $n \ll N$ the eigenbasis dimension, (3) *discretisation*: $x \in S''$ in domain 2^m , with m the bitstring length, and (4) *interpolation*: a in domain 2^m ; as illustrated below:

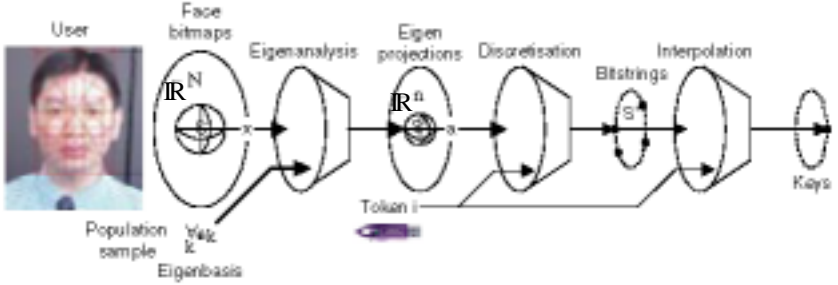


Fig. 1. Bio-hash representations and transformations

with enhanced stability at each step. Note this abstracted outlook does not take into account bitmap pre-processing prior to step (2), which is in actual fact extremely important due to the obvious correlation between the offset tolerances of (2) and (3). Enhancements in the former can be effected via application of Hambridge feature location [12] and eigenanalysis as reported in Ngo-Goh [13]. Our methodology is still straightforwardly applicable, with $\mathbf{a}$ and $\mathbf{x}$ in this case a concatenation of feature-specific contributions.

The primary concern from the security viewpoint centres on protection of information during the representational transformations, and in particular whether these transformations can be inverted to recover the input information. The above-listed parameters are said to be zero knowledge (ZK) representations of their inputs if the transformations are non-invertible, as in the case of cryptographic hash $h(i, j) : 2^m \times \forall_j 2^{m'} \rightarrow 2^m$ for token serialisation i and secret knowledge j . This

motivates an equivalent level of protection for biometric $\mathbf{a}$; which is accomplished via token-specification of the (3) and (4) representations, such that bio-hash $H(i, \mathbf{a}) : 2^m \times \mathbb{R}^n \rightarrow 2^m$ does not jeopardise $\langle i, \mathbf{a} \rangle$. ZK representation $a = H(i, \mathbf{a})$ is

subsequently useful for standard cryptographic operations ie signature generation and message decryption. Note H has an important (and challenging) additional requirement over h , namely offset tolerance so that $H(i, \mathbf{a})$ is stable for $\forall \mathbf{a} \in S'$. This requirement essentially addresses the fundamental gap between biometric similarity and cryptographic equality.

Our methodology is outlined in the above-discussed stages, as follows:-

3.1 Biometric Eigenanalysis

Sirovich-Kirby principal components analysis (PCA) presumes that $\mathbb{R}^N$ face bitmaps are more effectively represented as $\mathbb{R}^n$ eigenprojections, with interim dimensionality $M \ll N$ corresponding to the number of distinct users in the bitmap database. *Eigenface* characterisation requires computation of eigenbasis $\mathbf{e}_k$ (ranked by eigenvalue c_k significance) for $k = 1 \dots M$. The $n \ll M$ principal eigenfaces enables descriptive accuracy up to an externally specified accuracy, with user-specific data represented as $\forall \mathbf{a} \in \mathbb{R}^n$ are then computed as $\mathbf{a}_k = \mathbf{e}_k^\dagger \cdot \mathbf{d}_\alpha$.

Conventional biometrics requires storage of user-specific $\mathbf{a}$ so as to provide a reference against freshly captured test data. This is not satisfactory from the security viewpoint, as an intercepted $\mathbf{a}$ opens up the possibility of transaction fraud. Revocation of $\mathbf{a}$ (analogous to password refreshment or token replacement) is also highly problematic for all biometric forms, and impossible for face data. This dilemma is a major motivation for our work, particularly in its emphasis that stored references are fundamentally insecure and that bio-hashing should operate in a *one-way* manner on fresh data, analogous to password hashing.

3.2 Biometric Discretisation

The most offset tolerant transformation on face data $\mathbf{a} \in \mathbb{R}^n$ is reduction down to a single-bit. This is accomplished via:-

1. Compute $s(\mathbf{a}, \mathbf{b}) = \sum_k c_k (a_k b_k)$ with random normalised $\mathbf{b} \in \mathbb{R}^n$
2. Assign $b(s) = \begin{cases} 0: & s < \mu - \sigma \\ 1: & s > \mu + \sigma \\ \emptyset: & s \in [\mu - \sigma, \mu + \sigma] \end{cases}$

for empirical μ and σ , the former of which should theoretically vanish due to above specification of $\mathbf{a}$ relative to the population average. Extracted $b(\mathbf{a} \cdot \mathbf{b})$ is a broad measure of whether $\langle \mathbf{a}, \mathbf{b} \rangle$ are inline or opposed, with σ applied to exclude the perpendicular case. This exclusion mitigates against data capture uncertainties in $\mathbf{a}$, which might otherwise result in bit-inversion for numerically small s .

Repetition of this procedure to obtain multiple bits raises the issue of inter-bit correlations, which is addressed via orthonormal set $\beta = \{\mathbf{b}_k: k=1 \dots v\}$ with $v < n$.

Each bit $x_k = b(\mathbf{a} \cdot \mathbf{b}_k)$ is hence rendered independent of all others, so that legitimate (and unavoidable) variations in $\forall \mathbf{a} \in S'$ that invert x_k would not necessarily have the same effect on x_k' .

Inter-bit correlations and observations thereof are also important from the security viewpoint, the latter of which is prevented via cryptographic hashing of the concatenated bits. Indeterminate bits $x_k = \emptyset$ are handled via replacement of near-perpendicular $\mathbf{b}_k$ with alternative $\mathbf{b}'_k$, the net effect of which is bit-extraction via adjusted set $\beta - \bigvee_{k \in \perp} \mathbf{b}_k + \bigvee_{k \in \perp} \mathbf{b}'_k$. This reformulation is facilitated by the original stipulation on v , which allows up to $n-v$ replacements for unsuitable $\mathbf{b}_k$.

The proposed discretisation via repeated inner-products then proceeds as follows:

1. Generate random $\beta + \bigvee \mathbf{b}'$ for $k = 1 \dots v \dots n$
2. Orthonormalise $\beta + \bigvee \mathbf{b}'$ via Gram-Schmidt procedure
3. For each $k = 1 \dots v$:
 1. Compute $s_k = \mathbf{a} \cdot \mathbf{b}_k$
 2. While $s_k \in [\mu - \sigma, \mu + \sigma]$:-
 1. Get next unused $\mathbf{b}'$
 2. Reassign $\mathbf{b}_k = \mathbf{b}'$ in β
 3. Recompute s_k
 3. Assign $x_k = b(s_k)$
4. Concatenate $\alpha = \bigvee_k x_k$
5. Compute $x = h(\alpha)$

Note the easy adaptability to the previously discussed multi-feature biometrics, and also the inherent scalability (with respect the $\alpha \in 2^v$ bitlength) equivalent to the Monroe et al methodology. The experimental data in the next section is designed to address signal processing issues, hence the omission of step (5) there. Step (3.2) is critical for representational stability ie the confinement of $x(\mathbf{a})$ for $\forall \mathbf{a} \in S'$ to a small set S'' , so as to facilitate mapping down to a single cryptographic key. This requires the generic stability of random $x(\mathbf{a} \cdot \mathbf{b}_k)$; and is a fundamental motivation for the presented error correction at two stages, the first of which uses σ valuation to mitigate against continuous-valued uncertainties in $\mathbf{a}$. The second-stage addresses the discretisation of these uncertainties in $x(\mathbf{a}) \in 2^m$.

Recall the stipulation that $\mathbf{a}$ be protected equivalent to other cryptographic key-factors, which is accomplished via the use of tokenised cryptographic mechanisms—ie X9.17 pseudorandom generators [14] constructed from ciphers or hashes—in step (1). Resultant sequence $\beta(i)$ and output $x(i, \mathbf{a})$ are hence ZK representations of i , and consequently protective of $\mathbf{a}$ as subsequently outlined; which is reminiscent of the Soutar et al methodology. Note the effect of different token i' on the β sequence, resulting in $x(i, \mathbf{a}) \neq x(i', \mathbf{a})$ to a high degree of certainty. The proposed tokenised

discretisation can therefore be said to combine the best attributes of the Soutar et al and Monroe et al approaches.

3.3 Cryptographic Interpolation

The limited uncertainty of $x \in S''$ is addressed via Shamir secret-sharing; which uses modular polynomial $f(x) : \mathbb{Z}_q \rightarrow \mathbb{Z}_q$ for secret encoding $f(0) = a$, which is the $2^m \subset \mathbb{Z}_q$ cryptographic key in our context. In the simplest linear case, this allows

secret recovery via $a = \frac{x \cdot f(x')}{x - x'} + \frac{x' \cdot f(x)}{x' - x} \pmod{q}$ with $x = x(i, \mathbf{a})$ and $x' = h(i)$. Coordi-

nate pair $\langle x, f(x) \rangle$ constitutes a secret-share, any two of which can be combined to recover a . The operational concept is to match one of the biometric-associated shares with the token-associated one, so as to be consistent with the above-outlined discussion on token-specific discretisations. This is a rigorous 2-of- μ threshold system, with $\mu = |S''|$ the number of possible discretisation outcomes corresponding to a particular user.

The still approximate nature of the above-presented discretisation is addressed via prior specification and token-side insertion of $X = \langle \forall \langle \chi, y \rangle, c \rangle$, with: (1) $\chi = h(x)$

and $y = \frac{x' \cdot f(x)}{x' - x} \pmod{q}$ for $\forall \mathbf{a} \in S'$, and (2) $c = f(x') \pmod{q}$; corresponding to some

random key-encoding polynomial f . Key-computation then commences as follows:-

1. Retrieve X from token
2. Compute $x = x(i, \mathbf{a})$ as previously outlined
3. Select y such that $\chi = h(x)$ else stop
4. Compute $a = \frac{c \cdot x}{x - h(i)} + y \pmod{q}$

provided $\forall x \in S''$ have been properly identified. Note that $a(i, \mathbf{a})$ in step (4) cannot be computed without one of the correct discretisations x or token i , and that neither of these can be recovered from the ZK representations in X . The latter can in fact be stored completely in the open, which is illustrative of the protocol-level security comparable to the Davida et al and Monroe et al formulations. This is in complete contrast to the highly sensitive handling of biometric references, and the serious consequences arising from failure thereof. It is furthermore possible to encode a

password-associated key-share via prior specification in X of $y'' = \frac{x' \cdot f(x'')}{x' - x''} \pmod{q}$

with $x'' = h(j)$ from password j . This enables subsequent token+knowledge computa-

tion $a = \frac{c \cdot x''}{x'' - h(i)} + y'' \pmod{q}$ via y'' from token-side X as a backup option when the

usual token+biometric computation is inapplicable ie in low-light conditions.

Polynomial thresholding is rigorous and versatile, but is on the other hand restrictive in that (small) μ has to known a priori. This requires a high degree of error tolerance in the above-discussed discretisation $x(i, \mathbf{a})$, as would result from suitable ad-

justment of σ . Key-interpolation is interpreted as a final error-correcting step in this context, supplementing the basic robustness of random bit-extraction and the replacement of bits over-sensitive to legitimate variations in $\mathbf{a}$. End result $\mathbf{a} = H(i, \mathbf{a})$ is hence: (1) *sensitively* dependant on i : so that exact correctness is required for $\beta(i)$ and $x'(i)$, the former of which contributes sensitively towards $x(i, \mathbf{a}) \in S''$, (2) *robustly* dependant on $\mathbf{a}$; commensurate with the discrete i and continuous $\mathbf{a}$ key-factors.

4 Experimental Data

The proposed methodology is tested on Spacek's *Faces94* dataset [15] posted on the University of Essex Website. This dataset contains frontal face photos taken from a fixed camera distance, with the subjects asked to speak throughout the process; resulting in biometric data with the following characteristics: (1) *data-base size*: 153 individuals, 3060 images, (2) *bitmap dimension*: 180×200 pixels, 256-level grayscale, (3) *photo illumination*: relatively uniform, with dark background, (4) *face scale in image*: relatively uniform, (5) *face position in image*: minor variations, (6) *face aspect*: very minor variations in turn, tilt and slant, and (7) *face expression*: significant variation due to speech. Faces94 is considered to be somewhat less challenging in comparison to other widely analysed datasets (ie Faces95 and Faces96) from the viewpoint of scale, aspect and illumination off-sets; but is excellent for our purposes as it simulates our anticipated operational scenario ie individual users in desktop or kiosk environments. There scenarios allow biometric capture under relatively controlled conditions, with users safely presumed to be facing forward in adequately illuminated surroundings. Recall the focus of this paper on the effects of post-eigenanalytic discretisation and error correction; hence our omission of image-preprocessing, which is acceptable for Faces94 but far less so for the other datasets. We look forward to presenting a more comprehensive analysis—with more challenging data, and incorporating image-preprocessing—in a subsequent publication. Faces94 is furthermore quite large with 20 distinct images per person; so that half can be used for establishment of the population eigenbasis, and the rest for testing.

The featured experimental configurations are as follows: (1) *pca-n*: denoting $\mathbb{R}^n$ eigenanalysis, (2) *pca+d-n*: denoting 2^n $\sigma = 0$ discretisation without exclusion of *weak* inner-products, and (3) *pca+de-n*: denoting 2^n discretisation with σ error-correction based on analysis of inner-products computed from random and user-specific eigenprojections. The last configuration amounts to exclusion parameter σ amounting to selection of the n most significant inner-products from a random sample of size $n' > n$. This necessitates a $\mathbb{R}^{n'}$ eigenbasis, with $n'-n$ corresponding to the Hamming distances between same user discretisations. Our methodology requires relatively small Hamming distances in the *pca+d-n* configurations, which are then further reduced via error-correction for *pca+de-n*. We acquired experimental data for $n = 20, 30, 40, 50, 60, 70$ and 80 in all cases.

4.1 Same and Different User Histogrammes

Population-wide histogrammes for: (1) Euclidean distance between same and different user eigenprojections, (2) Hamming distance between same and different user discretisations; are presented below:-

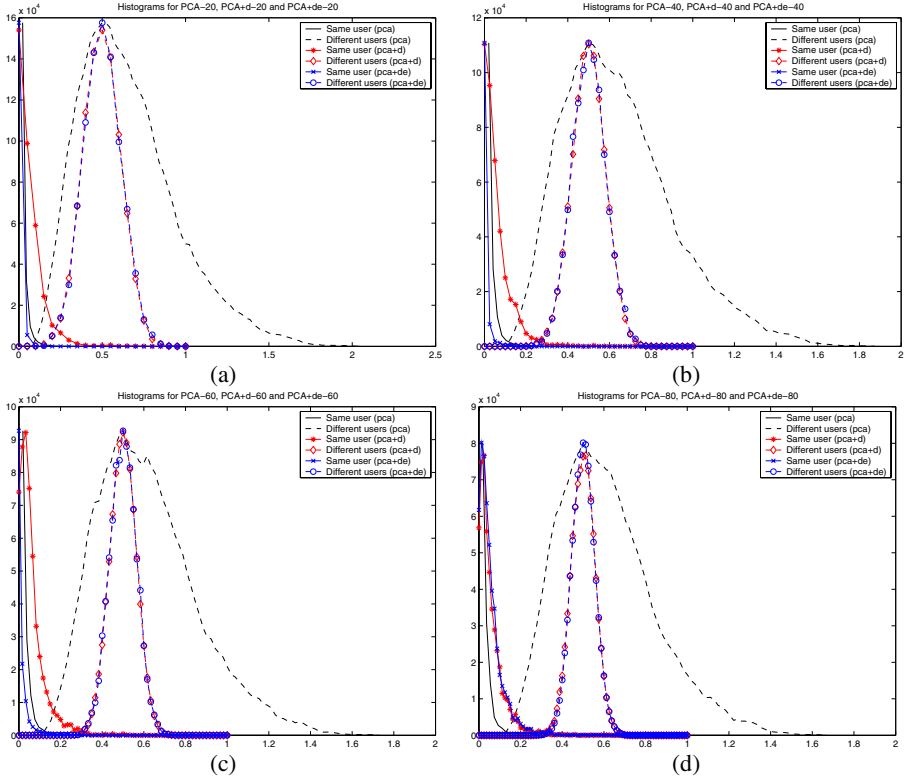


Fig. 2. Same and different user histogrammes for pca-n, pca+d-n and pca+de-n; for n = (a) 20, (b) 40, (c) 60 and (d) 80

on a normalised scale, with measure $\Delta x'$ derived from: (1) $\frac{\Delta x}{2x_p}$ for Euclidean dis-

tances, with x_p the peak of the different user histogrammes for pca-n, and (2) $\frac{\Delta x}{n}$ for Hamming distances arising from pca+d-n and pca+de-n.

Note the occurrence of histogramme peaks—for pca+d-n (red-highlighted) and pca+de-n (blue-highlighted)—at Hamming distances of 0 (same user) and $\frac{n}{2}$ (different users), both of which are strong vindications of the proposed methodology. Clear separation of the same and different user histogrammes is extremely important from the security viewpoint, hence the attractiveness of the pca+de-n same user histogrammes with its much steeper peak-to-plateau drop-offs compared to the corre-

sponding $\text{pca}+\text{d}-n$ profiles. The above-outlined Euclidean normalisation allows for qualitative comparison of $\text{pca}-n$ characteristics, which also emphasises the advantages of the $\text{pca}+\text{de}-n$ configurations.

These sharp drop-offs are clearly apparent in the $n = 40$ and 60 cases, but less so for $n = 20$ and 80 . This can be attributed to the descriptive insufficiency for low n , and over-sensitivity to noise for high n configurations; not just for the proposed $\alpha \in 2^n$ bitstrings but also for the basic. $\mathbf{a} \in \mathbb{R}^n$ The form of the $\text{pca}+\text{de}-n = 40$ and 60 histogrammes allows for specification of *zero* FAs without overly jeopardising the FR performance. FR (FA = 0) is, in fact, an important merit criteria in the proposed framework, which anticipates $H(i, \mathbf{a})$ parameterised cryptographic functionality. It is important to be able to preclude the occurrence of FAs in this context.

4.2 FA and FR Characteristics

Establishment of FR (FA = 0) and the more commonly cited crossover error (CE) rate (at which point FA = FR) for a particular configuration requires analysis of the FA-FR operational characteristics ie:

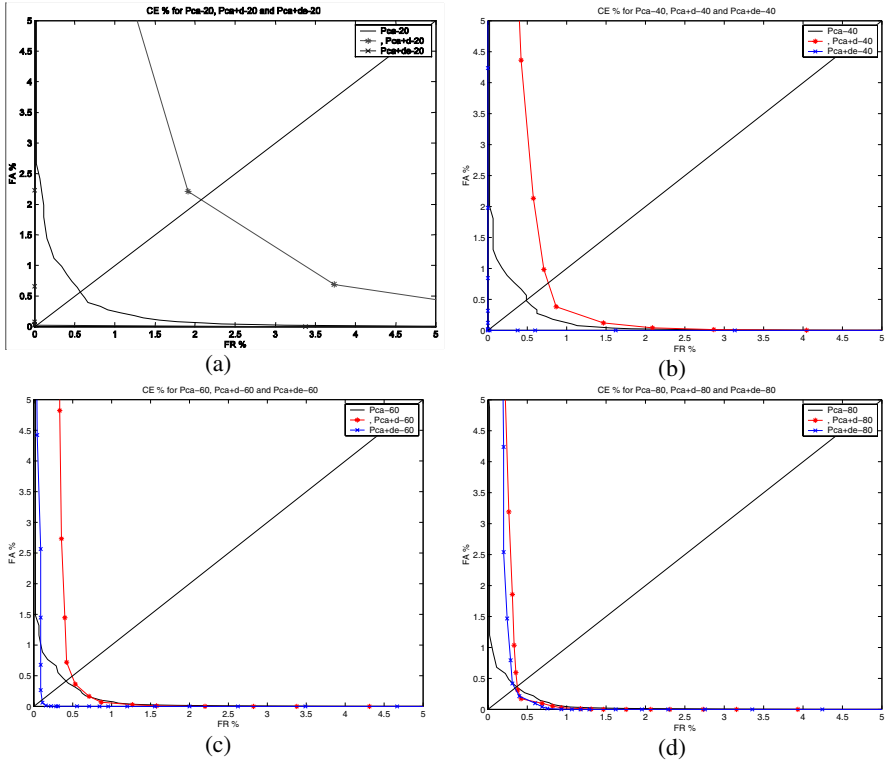


Fig. 3. Operational characteristics for for $\text{pca}-n$, $\text{pca}+\text{d}-n$ and $\text{pca}+\text{de}-n$; for $n =$ (a) 20, (b) 40, (c) 60 and (d) 80

Note the higher CE rates of *pca+d-n* (red-highlighted) compared to the corresponding *pca-n* configuration; the former of which eventually drops under the latter, corresponding to lower FR (FA = 0) rates. The *pca+d-n* configuration is hence more secure than the corresponding *pca-n*, but on the other hand somewhat less robust in terms of recognition. Error-correction can certainly be expected to improve recognition, as can be seen from the consistent location of *pca+de-n* (blue-highlighted) inside the corresponding *pca+d-n* profile. This reduces the CE point dramatically for the $n = 20, 40$ and 60 cases; but (in common with the Fig 3 histogrammes) less so for $n = 80$. The CE points for *pca+de-n* are in fact a significant improvement over the corresponding *pca-n*, again with the notable exception of the $n = 80$ case.

4.3 General Characteristics

The general characteristics of *pca-n*, *pca+d-n* and *pca+de-n* are as follows:

Table 1. Characteristics of (a) *pca-n*, (b) *pca+d-n* and *pca+de-n*

<i>a eigenbasis</i>	<i>Same user diff</i> (<i>Euclidean</i>)		<i>FR %</i> (<i>FA = 0</i>)		<i>CE %</i>	
20	0.030		4.47		0.56	
30	0.035		2.80		0.57	
40	0.041		2.49		0.49	
50	0.046		2.37		0.49	
60	0.049		2.26		0.41	
70	0.053		1.69		0.55	
80	0.055		1.60		0.37	

<i>α bitlength</i>	<i>Same user diff</i> (<i>Hamming</i>)		<i>FR %</i> (<i>FA = 0</i>)		<i>CE %</i>	
	<i>pca+d</i>	<i>pca+de</i>	<i>pca+d</i>	<i>pca+de</i>	<i>pca+d</i>	<i>pca+de</i>
20	1.15	0.03	29.70	3.37	2.07	0.02
30	1.72	0.09	8.42	0.01	1.02	0.01
40	2.18	0.15	4.04	0.01	0.77	0.01
50	2.85	0.47	1.98	0.27	0.57	0.07
60	3.53	0.86	1.57	0.22	0.49	0.10
70	4.17	1.79	1.17	0.35	0.55	0.15
80	4.48	4.51	1.31	0.93	0.37	0.34

and clearly illustrate the functional shortcomings of under and over-sized n representations. Choice of operational n in the (30, 45) range appears most suitable, so as to simultaneously avoid degraded recognition and frequent occurrence of bit-errors.

Note the relatively small Hamming distances between same user *pca+d-n* bit-strings, which vindicates the Section 3.2 discretisation and error-correction. This implies the sufficiency of relatively small $n'-n$ margins. The even smaller Hamming differences—less than a single-bit for most of the above-tabulated operational range—in the augmented *pca+de-n* case is also encouraging as it suggests a relatively small number of $x(i, \mathbf{a})$ outcomes per user, which is important for the Section 3.3 interpolation.

5 Security Analysis

The security of H should be evaluated in terms of key-factor: (1) *independence*: ie evaluation of $a = H(i, \mathbf{a})$ in the absence of i or $\mathbf{a}$, and (2) *non-recovery*: of i or $\mathbf{a}$ given specific value of $a(i, \mathbf{a})$ and the other factor; with the benchmark being cryptographic hashing of i and secret knowledge j . Recall that $a = h(i, j)$ cannot be computed without both $\langle i, j \rangle$ factors, so that adversarial deduction is no more probable than random guessing of order 2^{-m} . Factorisation $\langle i, j \rangle$ is also protected by the target-collision resistance of h , so that deduction of i or j —from output $a(i, j)$ and one of the factors—is equally improbable.

5.1 Key-Factor Independence

Non-possession of i means that tokenised $\beta(i)$ is unavailable to an adversary, so that previously intercepted (or fabricated) $\mathbf{a}$ is simply not useful. This prevents meaningful deduction of $a(i, \mathbf{a})$, with random guessing being of probability q^{-1} in this case. Possession of i is more useful as it divulges $\forall \chi = h(x)$ from the token-inserted $X(i, \mathbf{a})$, which suggests an analytic strategy whereby random $\alpha \in 2^v$ bitstrings are tested for suitability with respect condition $h(h(\alpha)) = \chi$. The collision probability is $\mu 2^{-v}$ in this case, hence the motivation to minimise μ and to maximise v . This is accomplished via suitable choice for inner-product exclusion parameter σ (which serves no useful purpose if over-large); and also by adoption of the previously discussed multi-feature eigenanalysis [13], so that lengthier α can be concatenated from feature-specific bitstrings. Note α with arbitrarily large v are straightforwardly obtained from integral transform representations, which do not restrict the length of the $\beta(i)$ sequence. Recall this issue of discretisation scalability also arises in the Monroe et al formulation.

The operational security of our scheme is enhanced via token-side access control and encryption of X , with respective parameterisation $\langle k, k' \rangle = h(i', i)$ for domain or platform serialisation i' . This necessitates prior token-side insertion of $\Psi = E_{k'}(X)$, with the following operational sequence:

1. Compute $\langle k, k' \rangle$ from token i
2. Transmit k to retrieve Ψ from token
3. Recover $X = D_k(\Psi)$

prior to the computations of Section 3, successful completion of which is restricted to domain/platform i' .

5.2 Key-Factor Non-recovery

Knowledge of $a(i, \mathbf{a})$ and $\mathbf{a}$ does not in any way jeopardise i , due to non-recovery of: (1) any $x \in S''$ from a , (2) any $\mathbf{b}_k \in \beta(i)$ from x and $\mathbf{a}$, and (3) i from $\beta(i)$ or any

subset thereof; thereby resulting in i deduction being no less probable than the 2^{-m} of random guessing. The other scenario of a and i compromise allows testing of random $a \in \mathbb{R}^n$ eigenprojections for suitability with respect condition $h(x(i, a)) = \chi$. Probability of a recovery in this case is μp^v , with $p < 2^{-1}$ due to exclusion of numerically small inner-products.

Key-factor protection is enhanced via reasonable operational measures: (1) minimisation of μ and maximisation of v , and (2) access control and encryption of X ; in addition to incorporation of i' dependence in the β sequence. This $\beta(i, i')$ specification is straightforwardly accomplished ie via initialisation $b_0(i')$ for the proposed X9.17 pseudorandom generator.

5.3 Cryptographic Applicability

The above-outlined $a = H(i, a)$ computation facilitates the application of asymmetric cryptographic protocols, ie for (1) online verification over a priori insecure environments, or (2) offline commitment and subsequent verification in relation to specific data; without presumptions that might be operationally inconvenient or unrealistic ie the establishment of communications security prior to biometric verification. Secure channel establishment in any case requires cryptographic support—ie the Diffie-Hellman (DH) [14] protocol—hence the motivation for the integrated handling of biometric and communications security, as subsequently outlined. Bio-hash H allows for cryptography predicated on $\langle i, a \rangle$ possession, which is more secure (due to simplicity of the key-computation conditions) and furthermore supportive of greater functional sophistication.

Cryptographic operations are straightforwardly parameterised via discrete logarithmic (DL) [14] or elliptic curve (EC) [16] key-pair of form $\langle a(i, a), A(a) \rangle$ with public-key $A = a \cdot g$ for basepoint g in some scalar-multiplicative subgroup $G_q \subseteq E$ of the specified curve. User-specific key-pair $\langle a, A \rangle$ is hence a ZK representation of $\langle i, a \rangle$ via the H and $g(a) : \mathbb{Z}_q \rightarrow G_q$ transformations, with remote identification in terms of $A(i, a) \in G_q$. This is qualitatively superior compared to the insecure and functionally limited $a \in \mathbb{R}^n$ of conventional biometrics.

6 Concluding Remarks

This paper outlines error-tolerant discretisation and cryptographic key-computation from user-specific face images and uniquely serialised tokens. Our bio-hash methodology has significant functional advantages over conventional biometrics ie extremely clean separation of the same and different user histogrammes and near-zero CE point, thereby allowing elimination of FAs without suffering from increased occurrence of FRs. $H(i, a)$ is furthermore highly secure with respect independence and non-

recovery of the $\langle i, \mathbf{a} \rangle$ key-factorisation, with tokenised immunity against biometric interception or fabrication. Use of token+biometric key $a(i, \mathbf{a})$ within the context of asymmetric cryptography is also attractive in that it enables secure and versatile functionality.

References

1. A Bodo (1994). *Method for Producing a Digital Signature with Aid of a Biometric Feature*. German Patent DE 42–43–908–A1
2. C Soutar & GJ Tomko (1996). *Secure Private Key Generation Using a Fingerprint*. Cardtech/Securetech Conf 1: pp 245–252
3. C Soutar, D Roberge, A Stoianov, R Gilroy & BVK Vijaya Kumar (1998). *Biometric Encryption Using Image Processing*. SPIE 3314: pp 178–188
4. GI Davida, Y Frankel & BJ Matt (1998). *On Enabling Secure Applications Through Off-Line Biometric Identification*. IEEE Symp on Security & Privacy: pp 148–157
5. GI Davida, Y Frankel, BJ Matt, & R Peralta (1999). *On the Relation of Error Correction and Cryptography to an Off-Line Biometric-Based Identification Scheme*. Wkshop Coding & Cryptography: Paris France
6. F Monrose, MK Reiter & S Wetzel (1999). *Password Hardening Based on Keystroke Dynamics*. 6–th ACM Conf on Comp & Comms Security: pp 73–82
7. F Monrose, MK Reiter, Q Li & S Wetzel (2001). *Cryptographic Key Generation from Voice*. IEEE Symp on Security & Privacy: pp 202–213
8. A Shamir (1979). *How to Share a Secret*. ACM Comms 22 (11): pp 612–613
9. T Matsumoto, H Matsumoto, K Yamada & H Hoshino (2002). *Impact of Artificial “Gummy” Fingers on Fingerprint Systems*. SPIE 4677: ??
10. L Sirovich & M Kirby (1987). *A Low-Dimensional Procedure for Characterisation of Human Faces*. J Optical Soc 4 (3): pp 519–524
11. M Turk & A Pentland (1991). *Face Recognition Using Eigenfaces*. IEEE Conf Comp Vision & Pattern Recognition: pp 586–591
12. J Hambridge (1926). *The Elements of Dynamic Symmetry*. Yale Univ Press, New Haven USA
13. DCL Ngo & A Goh (2003). *Facial Feature Extraction via Dynamic Symmetry Modelling for User Identification*. Pattern Pattern Recognition Letters.
14. AJ Menezes, P van Oorschot & S Vanstone (1996). *Handbook of Applied Cryptography*. CRC Press, Boca Raton USA.
15. L Spacek (2000). *Face Recognition Data*. <http://cswww.essex.ac.uk/allfaces/index.html>
16. AJ Menezes (1993). *Elliptic-Curve Public-Key Cryptosystems*. Kluwer Academic Press, Boston USA.