

МЕТОДИ ПОБУДОВИ ТА АНАЛІЗУ СТІЙКИХ НЕЧІТКИХ ЕКСТРАКТОРІВ ДЛЯ БІОМЕТРИЧНИХ СИСТЕМ

Р. В. Сковрон^{1,а}

¹ Навчально-науковий Фізико-технічний інститут

Анотація

У даній роботі розглядається проблема побудови стійких нечітких екстракторів для біометричних систем. Проведено та проаналізовано експерименти щодо впливу методів зменшення розмірностей у поєднанні з різними схемами квантування на геометрію простору ознак. Також розглядаються основні підходи до побудови багатофакторних нечітких екстракторів.

Отримані результати дозволяють оцінити та адаптувати компроміс між завадостійкістю і криптографічною стійкістю системи та можуть бути використані при побудові біометричних криптосистем на основі нечітких екстракторів.

Ключові слова: нечіткий екстрактор, методи квантування, методи зменшення розмірності

Вступ

Біометричні системи автентифікації набувають широкого застосування завдяки відсутності необхідності запам'ятовування секретів. Однак біометричні ознаки є сильно зашумленими, що унеможливає їх пряме використання у криптографічних протоколах. Для розв'язання цієї проблеми використовують нечіткі екстрактори, що дозволяють генерувати стабільний ключ із зашумленого джерела. Проте їхня практична реалізація часто стикається з проблемою високого рівня біометричного шуму, що призводить до послаблення криптографічної стійкості.

У доповіді розглядається проблема побудови стійких нечітких екстракторів для біометричних систем та досліджується вплив методів зменшення розмірності простору біометричних векторних ознак та методів квантування на геометрію простору ознак, ентропію та рівень хибного допуску. Також розглянуто основні підходи до побудови багатофакторних нечітких екстракторів.

1. Основні означення та теоретичні відомості

Для оцінки криптографічної стійкості класична ентропія Шеннона недостатня, оскільки вона описує лише середню невизначеність. Натомість для більш точної оцінки використовують мінімальну ентропію [1], яка описує найгірший випадок непередбачуваності.

Нехай $\mathcal{X}$ — розподіл визначений над множиною X , тоді мінімальна ентропія визначається як

$$H_{\infty}(\mathcal{X}) = -\log_2 \left(\max_{x \in X} \Pr[\mathcal{X} = x] \right).$$

У контексті нечіткого екстрактора використовують означення середньої мінімальної ентропії [1] для оцінки його стійкості. Нехай $\mathcal{X} \times \mathcal{Y}$ спільний розподіл випадкових величин над множинами X та Y , тоді середня мінімальна ентропія визначається як

$$\tilde{H}_{\infty}(X|Y) = -\log_2 \left(\mathbb{E}_{y \leftarrow Y} \left[\max_{x \in X} \Pr[X = x|Y = y] \right] \right).$$

Така міра використовується для обґрунтування стійкості нечіткого екстрактора [1], тому що вона описує найгірший випадок при наявності додаткової інформації Y .

Екстрактор ознак f — це функція, зазвичай нейронна мережа, що відображає вхідні біометричні дані великої вимірності у компактний векторний простір. Для зображень це відображення має вигляд

$$f : \mathbb{R}^{W \times H \times C} \rightarrow \mathbb{R}^n,$$

де W, H, C — ширина, висота та кількість каналів вхідного зображення, а n — розмірність вихідного вектора ознак (ембедингу).

2. Методи квантування

Біометричні векторні дані, отримані з екстрактора ознак, належать до простору дійсних чисел $\mathbb{R}^n$, однак використання неперервних векторних величин у криптографічних цілях є непрактичним, тому виникає необхідність їхньої дискретизації. Для розв'язання цієї проблеми застосовують методи квантування.

Нехай $\mathcal{V} \subset \mathbb{R}^n$ — простір неперервних векторів ознак, а Σ — скінченний алфавіт, тоді $q : \mathcal{V} \rightarrow \Sigma^n$, називається функцією квантування. Головною вимогою до такої функції є максимальне збереження просторової геометрії вихідних векторів.

Рівномірне квантування (EQ) [2]. Цей метод квантує компоненти вектора шляхом розбиття

^аromanskovronn@gmail.com

діапазону їх значень на r рівномірних інтервалів за допомогою порогових значень. Вони обираються так, щоб забезпечити однакову ймовірність потрапляння значення у кожен з інтервалів згідно з емпіричним розподілом даних. У бінарному випадку $\Sigma = \{0, 1\}$ метод зводиться до використання одного порогу, наприклад, середнього значення або медіани $b_i = \mathbb{1}[v_i \geq \mu_i]$. Головною перевагою методу є максимізація ентропії на кожен згенерований символ. Проте він має високу чутливість до шуму поблизу порогових значень та ігнорує кореляції між різними компонентами вектора.

Випадкова проекція (RP) [3]. У цьому методі квантування вектора v відбувається шляхом множення на випадкову матрицю $\Pi \in \mathbb{R}^{m \times n}$ зі стандартного нормального розподілу $\Pi_{i,j} \sim \mathcal{N}(0, 1)$. Квантування вектора $v \in \mathcal{V}$ обчислюється як

$$\hat{v} = \text{sign}(\Pi v).$$

Головною перевагою методу є його універсальність, він не залежить від вихідного розподілу даних, не потребує навчання, має доведені теоретичні гарантії збереження відстаней між векторами [3] та можливість змінювати розмірність вхідного вектора. Однак недоліками є ігнорування емпіричного розподілу даних вихідних ознак, а також висока обчислювальна та просторова складність $\mathcal{O}(m \cdot n)$.

Циркулянтний бінарний ембедінг (СВЕ) [4]. Є модифікацією попереднього підходу, замість повністю випадкової матриці використовується циркулянтна матриця $\mathbf{R} = \text{circ}(\mathbf{r}) \in \mathbb{R}^{n \times n}$, задана вектором $\mathbf{r} \in \mathbb{R}^n$ зі стандартного нормального розподілу $\mathcal{N}(0, 1)$, та діагональна матриця випадкових знаків $\mathbf{D} \in \mathbb{R}^{n \times n}$.

Квантування вектора $v \in \mathcal{V}$ обчислюється як $\hat{v} = \text{sign}(\mathbf{R}\mathbf{D}v)$. Проте завдяки циркулянтній структурі множення $\mathbf{R}\mathbf{x}$ еквівалентне циклічній згортці та може бути ефективно обчислене за допомогою дискретного перетворення Фур'є

$$\hat{v} = \text{sign}(\mathcal{F}^{-1}(\mathcal{F}(\mathbf{r}) \circ \mathcal{F}(\mathbf{D}v)))$$

де $\mathcal{F}$ — дискретне перетворення Фур'є, а $\circ$ — покомпонентний добуток Адамара.

Цей метод є оптимізацією методу випадкової проекції зі зниженням часової складності до $\mathcal{O}(n \log n)$ та просторової до $\mathcal{O}(n)$. Також існує можливість оптимізації генеруючого вектора матриці $\mathbf{R}$ для конкретного розподілу.

3. Метрики ефективності екстрактора ознак

Основними та найбільш інформативними метриками для оцінки ефективності екстрактора ознак f при фіксованому порозі прийняття рішення ε є коефіцієнт помилкової відмови FRR та коефіцієнт помилкового допуску FAR .

Нехай ρ_{same} — розподіл пар біометричних ознак однієї особи, а ρ_{diff} — для різних осіб, тоді ймовірності помилок визначаються як

$$FRR(\varepsilon) = \Pr_{(\mathbf{x}, \mathbf{x}') \sim \rho_{\text{same}}} [d(f(\mathbf{x}), f(\mathbf{x}')) > \varepsilon],$$

$$FAR(\varepsilon) = \Pr_{(\mathbf{x}, \mathbf{y}) \sim \rho_{\text{diff}}} [d(f(\mathbf{x}), f(\mathbf{y})) \leq \varepsilon],$$

де d — функція відстані. На практиці ці значення оцінюються емпірично як частка помилкових класифікацій на тестових вибірках $\mathcal{P}_{\text{same}}$ та $\mathcal{P}_{\text{diff}}$.

Показником загальної якості системи є рівень рівної помилки EER значення, при якому $FAR(\varepsilon) = FRR(\varepsilon)$. Чим нижче значення EER , тим кращою є роздільна здатність екстрактора ознак. Однак для збільшення безпеки поріг ε зазвичай зміщують у бік мінімізації FAR .

Сучасні нейронні моделі здатні досягати значень $FAR \approx 2^{-20}$ [5], що визначає верхню межу безпеки системи. Це значення можна інтерпретувати як ефективну мінімальну ентропію екстрактора ознак $H_{\infty} \approx -\log_2(FAR(\varepsilon))$ [6], відповідно, такі показники забезпечують близько 20–30 біт ентропії.

Для побудови нечіткого екстрактора та аналізу його стійкості важливе значення мають математичні сподівання внутрішньокласових d^+ та міжкласових d^- відстаней

$$d^+ = \mathbb{E}_{(\mathbf{x}, \mathbf{x}') \sim \rho_{\text{same}}} [d(f(\mathbf{x}), f(\mathbf{x}'))],$$

$$d^- = \mathbb{E}_{(\mathbf{x}, \mathbf{y}) \sim \rho_{\text{diff}}} [d(f(\mathbf{x}), f(\mathbf{y}))].$$

В ідеалі відносні відстані прямують до $d^+ \rightarrow 0$ та $d^- \rightarrow 0.5$. Проте на практиці, реальні значення d^+ зазвичай наближаються до 0.20–0.25. Цю різницю можна інтерпретувати як біометричний шум, компенсація якого є головним завданням механізмів корекції помилок у нечіткому екстракторі.

Якість розділення розподілів внутрішньокласових d^+ та міжкласових d^- відстаней можна оцінити за допомогою індексу роздільності d' (decidability index) [7], що обчислюється як

$$d' = \frac{|d^- - d^+|}{\sqrt{\frac{1}{2}(\sigma^{-2} + \sigma^{+2})}},$$

де σ^+ та σ^- — стандартні відхилення внутрішньокласових та міжкласових відстаней відповідно. Чим вище значення d' , тим меншим є перекриття розподілів ρ_{same} та ρ_{diff} , що безпосередньо гарантує нижче значення EER і вищу надійність системи.

4. Вплив методів зменшення розмірностей

Біометричні векторні ознаки, отримані з екстрактора ознак, зазвичай містять надлишкову інформацію та шум, що негативно впливає на ефективність нечітких екстракторів. Тому необхідно дослідити методи зменшення розмірності як потенційний засіб усунення біометричного шуму та оптимізації міжкласових і внутрішньокласових відстаней для покращення нечіткого екстрактора.

Для зменшення впливу шуму, стискання простору ознак та зниження обчислювальної складності застосовують методи зменшення розмірності, зокрема метод головних компонент (PCA), лінійний дискримінантний аналіз (LDA) [8] та метод випадкових проєкцій (RP) [3].

Окрім зменшення розмірності, такі перетворення змінюють просторову геометрію векторних ознак, що безпосередньо впливає на розподіли внутрішньокласових та міжкласових відстаней, значення FAR , а також ефективну мінімальну ентропію. Експериментальні результати, отримані з використанням моделі MagFace [5] на наборі даних LFW [9], наведено в таблиці 1. У всіх експериментах для квантування використовувався метод СВЕ.

Таблиця 1. Вплив методів зменшення розмірності на геометрію відстаней векторних ознак

Метод	dim	d^+	σ^+	d^-	σ^-	d'
Базовий	512	0.216	0.041	0.483	0.039	6.686
LDA	368	0.210	0.040	0.482	0.039	6.909
	256	0.211	0.042	0.483	0.047	6.168
	196	0.211	0.045	0.482	0.047	5.894
	128	0.167	0.046	0.483	0.062	5.841
	96	0.122	0.051	0.485	0.067	6.098
	64	0.113	0.053	0.475	0.084	5.166
PCA	368	0.236	0.046	0.498	0.038	6.187
	256	0.240	0.048	0.499	0.046	5.484
	196	0.243	0.048	0.498	0.046	5.399
	128	0.212	0.055	0.497	0.059	5.029
	96	0.204	0.059	0.499	0.067	4.676
	64	0.177	0.063	0.499	0.077	4.589
RP	368	0.231	0.056	0.490	0.062	4.395
	256	0.219	0.045	0.480	0.050	5.452
	196	0.212	0.049	0.483	0.052	5.367
	128	0.238	0.055	0.478	0.063	4.048
	96	0.204	0.058	0.481	0.072	4.231
	64	0.191	0.064	0.472	0.083	3.799

Було проведено оцінку мінімальної ентропії H_∞ при стисненні простору ознак до вимірності 256. Хоча зі стисненням простору загальна ентропія знижується, між методами спостерігаються відмінності у кількості збереженої H_∞ . Для методу PCA у поєднанні з алгоритмами квантування СВЕ та RP залишкова мінімальна ентропія становить ≈ 215 біт, що дає ≈ 0.84 біта на компоненту. Проте значна частина цього значення зумовлена високоентропійним внутрішньокласовим шумом, оскільки значення d^+ не демонструє відповідного зменшення. Метод рівномірного квантування показує найменшу втрату ентропії ≈ 0.94 біта на компоненту. Метод LDA та метод випадкових проєкцій показують нижчі значення залишкової ентропії ≈ 0.73 та ≈ 0.77 біта на компоненту відповідно.

Для оцінки FRR використовувалась така методика експерименту: для кожної конфігурації визначався поріг прийняття рішення ε за якого FRR становить $\approx 1\%$, після чого обчислювалося відповідне значення FAR .

Для методів RP та СВЕ значення порогу, необхідне для досягнення заданого рівня FRR , становить $\varepsilon \approx 0.13$, натомість для EQ поріг є вищим $\varepsilon \approx 0.21$. Відповідно, якби для методу EQ було застосовано такий самий низький поріг ε , як для СВЕ або RP, то показник FAR виявився б ще нижчим. Методи LDA та PCA у цьому контексті показали приблизно однакові результати, тоді як метод випадкових проєкцій виявився найменш чутливим до зміни розмірностей.

5. Завадостійкі коди

Для побудови нечітких екстракторів необхідний механізм компенсації природного біометричного шуму. Одним з таких способів є завадостійкі коди. Завадостійкий код можна визначити як механізм внесення надмірності у дані для подальшого виявлення та виправлення помилок.

Нехай Σ – скінченний алфавіт, завадостійкий код C називається (n, k, d) -кодом, якщо він є підмножиною Σ^n , призначеною для кодування слів з алфавіту Σ^k , а мінімальна відстань між різними кодовими словами дорівнює d . Якщо алфавіт заданий над скінченним полем $\Sigma = \mathbb{F}_q$, а C утворює лінійний підпростір $\mathbb{F}_q^n$ розмірності k , тоді такий код C називається q -арним лінійним кодом. Кодування інформаційного слова $m \in \Sigma^k$ у кодове слово визначається відображенням $\text{Encode}_C(m)$, а декодування зашумленого кодового слова відображенням — $\text{Decode}_C(c')$.

Бінарні завадостійкі коди найпрактичніші для нечітких екстракторів через випадковий характер біометричного шуму. Коди з ітеративним декодуванням, незважаючи на високу здатність до виправлення помилок, не забезпечують гарантованого виправлення фіксованої кількості помилок, тоді як алгебраїчні коди надають детерміновані гарантії декодування.

6. Нечіткі екстрактори

Класичні криптографічні механізми вимагають точного збігу ключів, що робить їх несумісними з біометричними даними, які за своєю природою є зашумленими. При кожному зчитуванні однієї й тієї ж біометричної ознаки отримані вектори відрізняються. Для розв'язання цієї проблеми використовуються нечіткі екстрактори (Fuzzy Extractors) [1] – криптографічні примітиви, що дозволяють генерувати ключ із зашумленого джерела і відтворювати його за умови достатньої близькості до оригіналу.

Нехай $(\mathcal{W}, d)$ – метричний простір. Тоді пара алгоритмів генерації та відновлення (Gen, Rep) називається $(\mathcal{W}, m, \ell, t, \epsilon)$ -нечітким екстрактором, якщо вона задовольняє такі умови:

1. $\text{Gen}(w) \rightarrow (R, P)$. Імовірнісний алгоритм, який приймає на вхід біометричний зразок $w \in \mathcal{W}$ та генерує приватний ключ $R \in \{0, 1\}^\ell$ та публічний допоміжний рядок $P \in \{0, 1\}^*$.
2. $\text{Rep}(w', P) \rightarrow R$. Детермінований алгоритм, який приймає зашумлений зразок $w' \in \mathcal{W}$ та допоміжний рядок P . Для будь-якого w' , що задовольняє умову $d(w, w') \leq t$, алгоритм відновлює початковий ключ R .
3. Для будь-якого розподілу W над $\mathcal{W}$ такого, що $H_\infty(W) \geq m$, нехай $w \leftarrow W$ та $(R, P) \leftarrow \text{Gen}(w)$. Тоді виконується $\Delta((R, P), (U_\ell, P)) \leq \epsilon_{\text{sec}}$, де U_ℓ — рівномірний розподіл над $\{0, 1\}^\ell$, а Δ — статистична відстань.

Для нечітких екстракторів існує декілька механізмів виправлення помилок, зокрема завадостійкі лінійні коди [1] та ймовірнісні алгоритми [10]. Класичною реалізацією нечіткого екстрактора є конструкція кодового зсуву [1]. Ідея полягає у маскуванні біометричного ембедінгу випадковим кодовим словом.

Нехай C лінійний $[n, k, d]$ -код, а Ext строгий екстрактор випадковості [1], тоді пара алгоритмів (Gen, Rep) називається конструкцією кодового зсуву, якщо:

Gen :	Rep :
1. $\underline{\text{Bxid}}: w \in \{0, 1\}^n$.	1. $\underline{\text{Bxid}}: (w', s)$.
2. $x \xleftarrow{\$} \{0, 1\}^k$.	2. $c' := w' \oplus s$.
3. $c \leftarrow \text{Encode}_C(x)$.	3. $c \leftarrow \text{Decode}_C(c')$.
4. $s := w \oplus c$.	4. $w := s \oplus c$.
5. $R \leftarrow \text{Ext}(w)$.	5. $R \leftarrow \text{Ext}(w)$.
6. $\underline{\text{Buxid}}: (R, s)$.	6. $\underline{\text{Buxid}}: R$.

Якщо виконується умова $d(w, w') \leq t$, то алгоритм Decode_C здатен відновити початкове кодове слово c та відповідно відновити секрет w . Однак використовувати вектор w як приватний ключ недоцільно, через відсутність достатньої мінімальної ентропії. Тому на останньому кроці застосовується строгий екстрактор випадковості [1], найчастіше для цього використовується криптографічна хеш-функція.

7. Проблематика нечітких екстракторів

При побудові нечіткого екстрактора існує фундаментальний компроміс між завадостійкістю та безпекою. Наявний високий природний рівень біометричного шуму на пряму впливає на необхідну здатність коду до виправлення помилок, що, у свою чергу, впливає на витік ентропії через публічні допоміжні дані $\tilde{H}_\infty(W | P) \leq \tilde{H}_\infty(W) - (n - k)$. Окрім цього, екстрактор ознак може бути недостатньо безпечним для використання, оскільки його гарантії безпеки зводяться до $\text{FAR}(\varepsilon)$. Відповідно, стійкість нечіткого екстрактора зводиться до $\mathcal{O}(\min(2^{\tilde{H}_\infty(W|P)}, 1/\text{FAR}(\varepsilon)))$. Однак на цьому проблеми не закінчуються, існують проблеми багаторазового використання, зв'язуваності біометрії, активних атак та інші [10, 11].

8. Методи багатофакторної агрегації

Використання єдиного джерела біометричних даних зазвичай не здатне забезпечити необхідний рівень криптографічної стійкості ключа. Багатофакторні нечіткі екстрактори є підходом до вирішення цієї проблеми шляхом агрегації кількох джерел ентропії [12, 13, 14]. Їх можна поділити на дві категорії: зашумлені та стабільні фактори.

Нехай (M, d) — метричний простір, тоді джерело ентропії з розподілом Φ_U називається стабільним, якщо $\mathbb{E}_{w, w' \leftarrow \Phi_U}[d(w, w')] = d^+ = 0$, інакше джерело називається зашумленим. Стабільні фактори є детермінованими джерелами, до них належать паролі, PIN-коди, криптографічні токени тощо. Зашумленими факторами є ті, що демонструють варіативність при повторному зчитуванні, наприклад біометрія. Підходи до агрегації факторів залежать від етапу обробки факторів та архітектури системи.

Агрегація на рівні ознак. Вектори ознак w_i конкатенуються у спільний вектор $w_{\text{agg}} = w_1 \parallel \dots \parallel w_n$, до якого застосовується нечіткий екстрактор

$\text{Gen}(w_{\text{agg}})$. Хоча теоретично початкова ентропія зростає адитивно $H_\infty(w_{\text{agg}}) = \sum H_\infty(w_i)$, також зростає загальна кількість помилок $t_{\text{agg}} = \sum t_i$. А якщо серед компонент w_i є різні вимірювання одного й того ж біометричного джерела, то їхня спільна мінімальна ентропія практично не збільшується, оскільки базовий секрет залишається незмінним, однак сумарний шум t_{agg} зростає пропорційно кількості вимірювань.

Окремим випадком цього підходу є додавання стабільних факторів, їхнє використання дозволяє зменшити відносну частку помилок. Однак такий підхід не завжди працює, оскільки через властивості лінійності кодів, стабільні фактори можуть бути ізольовані від зашумлених та атаковані окремо повним перебором [15].

Каскадні перестановки. Стабільний фактор у такому підході використовується як секретний параметр для ініціалізації псевдовипадкової перестановки або ортогонального перетворення зашумленого вектора $\tilde{w} = \pi_s(w)$ до чи після квантування, який потім подається на вхід нечіткого екстрактора $\text{Gen}(\tilde{w})$. Цей підхід тісно пов'язаний з ідеєю скасовної біометрії (cancelable biometrics) [16], де до вектора ознак застосовуються незворотні перетворення, що зберігають геометрію відстаней.

Оскільки такі перетворення π_s не впливають на геометрію відстаней, то і розподіл внутрішньокласової відстані d^+ зберігається, а тому не вимагає зміни параметрів завадостійкого коду для легітимного користувача. Натомість спроба автентифікації з хибним стабільним фактором змінює структуру вектора, збільшуючи кількість помилок вище здатності коду до їхнього виправлення [17, 11]. Також зберігається властивість скасовної біометрії — можливості відкличкання та оновлення стабільного фактора.

Теоретично, такий підхід має гарантувати мультиплікативну складність атаки $\mathcal{O}(2^{H_\infty(S)} \cdot 2^{\tilde{H}_\infty(W|P)})$, шляхом унеможливлення ізольованої компрометації джерел ентропії. Проте на практиці, загальна стійкість системи критично залежить від ентропії стабільного фактора S . Через нерівномірність розподілу біометричних ознак, зломисник може здійснити статистичну розрізняльну атаку. Перебираючи простір стабільного фактора, зломисник може перевіряти результат $\pi_s^{-1}(w)$ на відповідність природному розподілу вектора W , що потенційно дозволить статистично виокремити правильний фактор S і звести складність атаки до адитивної.

Послідовне криптографічне зв'язування. Цей підхід встановлює строгу ієрархію між факторами шляхом каскадного шифрування публічних допоміжних даних [15]. Ключ R_1 , отриманий зі стабільного фактора використовується для шифрування допоміжного рядка зашумленого фактора $P_2^* = \text{Enc}(R_1, P_2)$. Одним з прикладів реалізації такого підходу є нечіткий екстрактор, захищений паролем [12, 13]. Це забезпечує семантичну безпеку допоміжного рядка P_2 до моменту компрометації R_1 , гарантуючи мультиплікативну складність атаки $\mathcal{O}(2^{H_\infty(S)} \cdot 2^{\tilde{H}_\infty(W|P_2)})$.

Загалом, послідовне криптографічне зв'язування є одним із сучасних методів агрегації. Альтернативним вектором розвитку багатофакторних нечітких

екстракторів є імовірнісний підхід на основі ідеї цифрового замка [10].

Висновки

У роботі досліджено вплив методів зменшення розмірності на біометричні ознаки. Показано, що у поєднанні з методами квантування вони суттєво впливають на геометрію простору ознак та *FAR*.

Метод LDA продемонстрував найкращі результати щодо зменшення внутрішньокласового шуму та придатність як етапу попередньої обробки даних для нечітких екстракторів. Рівноймовірне квантування разом із PCA та LDA дозволяє покращити показник *FAR*, зменшуючи біометричний шум та зберігаючи високий рівень мінімальної ентропії. Отримані результати дозволяють точніше визначити та адаптувати компроміс між завадостійкістю і криптографічною стійкістю при побудові біометричних систем на основі багатofакторних нечітких екстракторів.

Перелік використаних джерел

1. Fuzzy Extractors: How to Generate Strong Keys from Biometrics and Other Noisy Data / Y. Dodis, R. Ostrovsky, L. Reyzin, A. Smith // *SIAM Journal on Computing*. — 2008. — Vol. 38, no. 1. — P. 97–139. — DOI: [10.1137/060651380](https://doi.org/10.1137/060651380).
2. WiFaKey: Generating Cryptographic Keys from Face in the Wild / X. Dong, H. Zhang, Y. L. Lai, Z. Jin, J. Huang, W. Kang, A. B. J. Teoh // *IEEE Transactions on Instrumentation and Measurement*. — 2024.
3. Yi X., Caramanis C., Price E. Binary embedding: Fundamental limits and fast algorithm // *International Conference on Machine Learning*. — PMLR. 2015. — P. 2162–2170.
4. On Binary Embedding using Circulant Matrices / F. X. Yu, A. Bhaskara, S. Kumar, Y. Gong, S.-F. Chang. — 2015. — arXiv: [1511.06480](https://arxiv.org/abs/1511.06480) [cs.DS]. — URL: <https://arxiv.org/abs/1511.06480>.
5. MagFace: A Universal Representation for Face Recognition and Quality Assessment / Q. Meng, S. Zhao, Z. Huang, F. Zhou. — 2021. — arXiv: [2103.06627](https://arxiv.org/abs/2103.06627) [cs.CV]. — URL: <https://arxiv.org/abs/2103.06627>.
6. Unforgettable Fuzzy Extractor: Practical Construction and Security Model / O. Kurbatov, D. Zakharov, L. Antadze, V. Mashtalyar, R. Skovron, V. Dubinin. — 2025. — URL: <https://eprint.iacr.org/2025/1799>. Cryptology ePrint Archive, Paper 2025/1799.
7. Daugman J. How iris recognition works // *IEEE Transactions on Circuits and Systems for Video Technology*. — 2004. — Jan. — Vol. 14, no. 1. — P. 21–30. — ISSN 1558-2205. — DOI: [10.1109/TCSVT.2003.818350](https://doi.org/10.1109/TCSVT.2003.818350).
8. Prince S. J., Elder J. H. Probabilistic Linear Discriminant Analysis for Inferences About Identity // *2007 IEEE 11th International Conference on Computer Vision*. — 2007. — P. 1–8. — DOI: [10.1109/ICCV.2007.4409052](https://doi.org/10.1109/ICCV.2007.4409052).
9. Labeled Faces in the Wild: A Database for Studying Face Recognition in Unconstrained Environments : tech. rep. / G. B. Huang, M. Ramesh, T. Berg, E. Learned-Miller ; University of Massachusetts, Amherst. — 10/2007. — No. 07–49.
10. Reusable Fuzzy Extractors for Low-Entropy Distributions / R. Canetti, B. Fuller, O. Paneth, L. Reyzin, A. Smith. — 2014. — URL: <https://eprint.iacr.org/2014/243>. Cryptology ePrint Archive, Paper 2014/243.
11. Ouda O., Nandakumar K., Ross A. Cancelable Biometrics Vault: A Secure Key-Binding Biometric Cryptosystem based on Chaffing and Winnowing // *2020 25th International Conference on Pattern Recognition (ICPR)*. — 01/2021. — P. 8735–8742. — DOI: [10.1109/ICPR48806.2021.9412957](https://doi.org/10.1109/ICPR48806.2021.9412957).
12. Robust and Reusable Fuzzy Extractor for Low-Entropy Distributions and Application to User Authentication / S. Panja, N. Tripathi, S. Jiang, R. Safavi-Naini // *IEEE Access*. — 2026. — Vol. 14. — P. 38230–38250. — DOI: [10.1109/ACCESS.2026.3670084](https://doi.org/10.1109/ACCESS.2026.3670084).
13. Tran H. Y., Hu J., Hu W. Biometrics-Based Authenticated Key Exchange With Multi-Factor Fuzzy Extractor // *IEEE Transactions on Information Forensics and Security*. — 2024. — Vol. 19. — P. 9344–9358. — DOI: [10.1109/TIFS.2024.3468624](https://doi.org/10.1109/TIFS.2024.3468624).
14. Nagar A., Nandakumar K., Jain A. K. Multibiometric cryptosystems based on feature-level fusion // *IEEE Transactions on Information Forensics and Security*. — 2012. — Vol. 7, no. 1. — P. 255–268.
15. Fang C., Li Q., Chang E.-C. Secure sketch for multiple secrets // *Proceedings of the 8th International Conference on Applied Cryptography and Network Security*. — Beijing, China : Springer-Verlag, 2010. — P. 367–383. — (ACNS'10). — ISBN 3642137075.
16. Patel V. M., Ratha N. K., Chellappa R. Cancelable Biometrics: A review // *IEEE Signal Processing Magazine*. — 2015. — Vol. 32, no. 5. — P. 54–65. — DOI: [10.1109/MSP.2015.2434151](https://doi.org/10.1109/MSP.2015.2434151).
17. Fuzzy commitment scheme for generation of cryptographic keys based on iris biometrics / S. Adamovic, M. Milosavljevic, M. Veinovic, M. Sarac, A. Jevremovic // *IET Biometrics*. — 2017. — Vol. 6, no. 2. — P. 89–96. — DOI: <https://doi.org/10.1049/iet-bmt.2016.0061>. — eprint: <https://ietresearch.onlinelibrary.wiley.com/doi/pdf/10.1049/iet-bmt.2016.0061>. — URL: <https://ietresearch.onlinelibrary.wiley.com/doi/abs/10.1049/iet-bmt.2016.0061>.