



METHODS OF SIZE REDUCTION TO INCREASE THE EFFICIENCY OF PERSONAL IDENTIFICATION BY VOICE

Nurimov Parakhat Baymuratovich

“Tashkent Institute of Irrigation and Agricultural Mechanization
Engineers” MTU

paranur87@gmail.com

<https://doi.org/10.5281/zenodo.17085406>

ARTICLE INFO

Received: 01st September 2025

Accepted: 08th September 2025

Online: 09th September 2025

KEYWORDS

Speaker recognition, dimensionality reduction, PCA, ICA, LDA, Genetic Algorithm, resource-constrained systems, real-time applications.

ABSTRACT

This paper addresses the problem of efficient speaker recognition on resource-constrained devices. The focus is placed on reducing memory and computational costs while preserving the discriminative power of the feature set. To achieve this, dimensionality reduction techniques such as Principal Component Analysis (PCA), Independent Component Analysis (ICA), Linear Discriminant Analysis (LDA), and the Genetic Algorithm (GA) were applied. Experimental results demonstrate that these approaches significantly reduce memory usage and computational complexity while maintaining high recognition accuracy. The proposed methodology is particularly suitable for mobile devices and real-time systems with limited resources.

SHAXSNI OVOZI ORQALI IDENTIFIKATSIYALASH SAMARADORLIGINI OSHIRISH UCHUN O'LCHAMNI QISQARTIRISH USULLARI

Nurimov Paraxat Baymuratovich

“Toshkent irrigatsiya va qishloq xo'jaligini mexanizatsiyalash
muhandislari instituti” MTU

paranur87@gmail.com

<https://doi.org/10.5281/zenodo.17085406>

ARTICLE INFO

Received: 01st September 2025

Accepted: 08th September 2025

Online: 09th September 2025

KEYWORDS

Shaxsni ovozi orqali identifikatsiyalash, o'lchamni qisqartirish, PCA, ICA, LDA, Genetik algoritmi, resurs cheklangan tizimlar, real vaqt ilovalari.

ABSTRACT

Ushbu maqolada cheklangan hisoblash resurslariga ega qurilmalarda shaxsni ovozi orqali tanib olish tizimini samarali tashkil etish masalalari ko'rib chiqiladi. Asosiy e'tibor xotira va hisoblash resurslarini kamaytirish, shu bilan birga belgilar to'plamlarining farqlash qobiliyatini saqlab qolishga qaratilgan. Shu maqsadda o'lchamni qisqartirish usullari – Asosiy komponentlar tahlili (PCA), Mustaqil komponentlar tahlili (ICA), Chiziqli diskriminant tahlil (LDA) va Genetik algoritmi (GA) qo'llanilgan. Tadqiqot natijalari ushbu



yondashuvlar yordamida xotira sarfi va hisoblash murakkabligi sezilarli darajada kamayishi, shu bilan birga aniqlik ko'rsatkichlari yuqori darajada saqlanishi mumkinligini ko'rsatadi. Ushbu yondashuv resurslari cheklangan mobil qurilmalar va real vaqt tizimlari uchun mos yechim sifatida qaralishi mumkin.

KIRISH

Shaxsni ovozi asosida tanib olish — biometrik xavfsizlik, nutq orqali autentifikatsiya va inson-mashina o'zaro ta'siri sohalarida dolzarb muammolardan biridir. Zamonaviy chuqur o'rganishga asoslangan usullar yuqori aniqlikka ega bo'lsa-da, ular katta hajmli hisoblash resurslari va grafik protsessorlarni talab qiladi. Bu esa hisoblash quvvati cheklangan qurilmalarda (mobil telefonlar, IoT qurilmalar, aqlli uy tizimlari va edge-devicelar) amaliy qo'llashni qiyinlashtiradi.

Klassik yondashuvlardan biri bo'lgan Gaussian Mixture Model – Universal Background Model (GMM-UBM) esa hisoblash nuqtayi nazaridan yengilroq bo'lib, resurslar cheklangan tizimlarda samarali ishlashi mumkin. Shuningdek, ovoz signalidan olinadigan belgilar to'plamlar ham modelning samaradorligiga bevosita ta'sir ko'rsatadi. Odatda, Mel Frequency Cepstral Coefficients (MFCC) va uning birinchi hamda ikkinchi tartibli hosilalari (Δ , $\Delta\Delta$) keng qo'llaniladi. Bunga qo'shimcha ravishda, Mel-spektral koeffitsiyentlar ham signalning spektral tuzilishini yanada to'liqroq ifodalaydi.

Mazkur maqolada MFCC, Δ , $\Delta\Delta$ va Mel-spektral belgilardan foydalanib, GMM-UBM asosida shaxsni ovozi orqali identifikatsiyalash masalasi ko'rib chiqildi. Shu bilan birga, belgilarni kamaytirish (Feature Reduction) usullarining samaradorligi tahlil qilinib, hisoblash quvvati cheklangan qurilmalarda amaliy qo'llash imkoniyatlari baholandi.

METODOLOGIYA

1. Ma'lumotlar to'plami

Tajriba uchun o'zbek tilida yozib olingan 100 ta shaxsdan iborat nutq ma'lumotlaridan foydalanildi. Har bir shaxs turli matnlarni talaffuz qildi, bu esa modelning matndan mustaqil (text-independent) shaxsni ovozi orqali identifikatsiyalashiga imkon yaratdi.

Audio fayllar 16 kHz chastotada yozib olindi va maksimal uzunlik 5 soniyaga normallashtirildi. Agar audio qisqaroq bo'lsa, nol qiymatlar bilan (padding) to'ldirildi.

Bu yondashuv o'zbek tilida shaxs tanish tizimlarini yaratishda amaliy imkoniyatlarni kengaytiradi, chunki o'zbek tilidagi ma'lumotlar bazalari xalqaro miqyosda kamroq o'rganilgan.

2. Belgilar to'plamini ajratish (Feature Extraction)

Shaxsni ovozi orqali identifikatsiyalash masalasi uchun quyidagi akustik belgilar to'plamlari olindi:

- MFCC (Mel Frequency Cepstral Coefficients) – asosiy 13 ta koeffitsiyent,
- Δ -MFCC (Delta) – vaqt bo'yicha birinchi tartibli hosila,
- $\Delta\Delta$ -MFCC (Delta-Delta) – ikkinchi tartibli hosila,



- Mel-Spektral koeffitsiyentlar – signalning chastotaviy taqsimotini ifodalaydi. Bu belgilar birlashtirilib yagona vektor shaklida saqlandi.

3. Tasniflash

Asosiy tasniflovchi sifatida Gaussian Mixture Model – Universal Background Model (GMM-UBM) ishlatildi. Bu model har bir shaxs uchun ehtimollik taqsimotini shakllantirib, test paytida shaxsning moslik darajasini baholaydi.

4. Belgilar to'plamini kamaytirish usullari

Resurslari cheklangan qurilmalarda ovoz asosida shaxsni tanib olish samaradorligini oshirish uchun o'lchamlarni kamaytirish (dimensionality reduction) metodlari qo'llanildi. Ushbu usullar eng muhim belgilarni tanlab olish, xotira va hisoblash xarajatlarini kamaytirish bilan birga, belgilar to'plamining farqlash qobiliyatini saqlab qolishga yordam beradi. Quyidagi metodlardan foydalanildi:

Asosiy komponentlar tahlili (PCA). PCA — bu statistik usul bo'lib, yuqori o'lchamli ma'lumotlarni maksimal dispersiyani saqlagan holda pastroq o'lchamli fazoga proyeksiyalaydi. PCA belgilar fazosida ortogonal yo'nalishlarni (asosiy komponentlarni) aniqlaydi va ma'lumotlarni shu yo'nalishlarga proyeksiyalaydi. Bu yondashuv ortiqcha ma'lumotlarni kamaytiradi va hisoblash jihatdan samarali bo'lgani uchun real vaqtli qo'llanmalarga mos keladi.

Mustaqil komponentlar tahlili (ICA). ICA belgilar fazosini statistik jihatdan mustaqil komponentlarga ajratadi. PCA dispersiyani maksimallashtirishga intilsa, ICA mustaqil bo'lgan no-Gaussiy signallarni ajratishga harakat qiladi. ICA aralashgan signallarni (masalan, turli nutq manbalarini) ajratishda foydali bo'lib, ba'zi hollarda shaxsga xos nozik belgilar to'plamlarini ushlashga yordam beradi.

Chiziqli diskriminant tahlil (LDA). LDA — bu nazoratli usul bo'lib, sinflararo farqlanishni maksimal darajada oshirishga qaratilgan. U belgilar vektorlarini shunday pastroq o'lchamli fazoga proyeksiyalaydiki, unda sinflararo dispersiya va sinf ichidagi dispersiya nisbatining qiymati maksimal bo'ladi. Shu tufayli LDA shaxsni ovozi orqali identifikatsiyalashsi va verifikatsiyasi kabi klassifikatsiya vazifalarida yuqori samaradorlik beradi.

Genetik algoritim (GA). GA — bu tabiatdagi evolyutsiya jarayonidan ilhomlangan optimizatsiyaga asoslangan belgilarni tanlash usuli. GA belgilar to'plamidan eng maqbulini tanlash uchun seleksiya, krossover va mutatsiya kabi genetik operatorlardan foydalanadi. Fitnes funksiyasi sifatida odatda klassifikatsiya aniqligi olinadi. GA o'qitish bosqichida hisoblash jihatdan ko'proq resurs talab qilsa-da, yakunda eng samarali belgilarni ajratib, yengilroq va tezkor ishlaydigan model yaratishga yordam beradi.

5. Baholash mezonlari (Evaluation Metrics)

Model samaradorligini baholash uchun quyidagi ko'rsatkichlardan foydalanildi:

- Aniqlik (Accuracy) – umumiy to'g'ri tasniflangan namunalar ulushi.
- F1-makro (F1-macro score) – sinflar bo'yicha muvozanatlangan ko'rsatkich.
- Top-k aniqlik (Top-1, Top-3, Top-5 Accuracy) – to'g'ri javobning yuqori k nomzod ichida joylashishini baholash.

NATIJALAR



Tajriba davomida GMM-UBM modeli turli belgilar to'plamlarini kamaytirish usullari bilan sinovdan o'tkazildi. Quyida 1-jadvalda MFCC, Δ , $\Delta\Delta$ va Mel-spektral belgilardan foydalanib olingan natijalar keltirilgan.

1-jadval. Belgilar to'plamini kamaytirish usullari bo'yicha identifikatsiyalash natijalari

Usul	Accuracy	F1-macro	Top-1	Top-3	Top-5
NONE	0.6742	0.6771	0.6742	0.7854	0.8333
PCA(40)	0.6364	0.6380	0.6364	0.7500	0.7854
ICA(40)	0.5960	0.5984	0.5960	0.7096	0.7702
LDA(79)	0.7753	0.7871	0.7753	0.8763	0.9141
GA(45)	0.7096	0.7068	0.7096	0.8232	0.8611

Natijalardan ko'rinib turibdiki:

- Hech qanday belgilar to'plamini kamaytirmasdan (NONE) ishlaganda aniqlik 67.42% bo'ldi.
- PCA va ICA natijani biroz pasaytirib yubordi, bu ularning shaxs belgilar to'plamlarini to'liq ifodalay olmaganidan dalolat beradi.
- GA (Genetik Algoritm) usuli 70.96% aniqlik ko'rsatdi va Top-5 aniqlikda 86.11% ga erishildi.
- Eng yaxshi natija LDA orqali kuzatildi: aniqlik 77.53%, F1-makro 78.71% bo'lib, Top-5 aniqlik 91.41% ga yetdi.

Bu shuni ko'rsatadiki, LDA belgilar to'plamini kamaytirishda eng samarali usul bo'lib chiqdi va GMM-UBM modelining umumiy ishlashini yaxshiladi. Shu sababli LDA asosida qurilgan model hisoblash resurslari cheklangan tizimlarda ham yuqori aniqlikka erishishi mumkin.

MUNOZARA

O'tkazilgan tajribalar shuni ko'rsatdiki, turli belgilar to'plamini kamaytirish usullaridan foydalanish GMM-UBM modelining samaradorligiga sezilarli ta'sir ko'rsatadi. MFCC, Δ , $\Delta\Delta$ va Mel-spektral belgilar kombinatsiyasi shaxsni ovozi orqali identifikatsiyalash uchun barqaror natija berdi.

Natijalar tahliliga ko'ra, PCA va ICA usullari belgilar to'plamini qisqartirishda informatsiya yo'qotishiga olib kelgani sababli aniqlik pasaydi. Biroq, Genetik Algoritm (GA) nisbatan yuqori natija berdi, chunki u evolyutsion optimallashtirish orqali eng mos belgilarni tanlash imkonini berdi.

Eng yaxshi natija LDA (Linear Discriminant Analysis) orqali olindi. LDA sinflararo chegaralarni hisobga olgan holda belgilar to'plamini kamaytirgani tufayli, shaxslarning akustik belgilar to'plamlarini samaraliroq ajrata oldi. Bu esa aniqlik va F1-makro ko'rsatkichlarida sezilarli yaqshilashga olib keldi. Shuningdek, Top-3 va Top-5 aniqlik darajalari ham yuqori bo'lib, tizim amaliy sharoitlarda bir nechta nomzod ichidan to'g'ri shaxsni tanib olish ehtimolini oshirdi.

Bu natijalar shuni ko'rsatadiki, klassik GMM-UBM yondashuvi chuqur neyron tarmoqlar darajasida bo'lmasa-da, resurslar cheklangan qurilmalarda qo'llash uchun amaliy yechim hisoblanadi. Ayniqsa, LDA kabi belgilarni kamaytirish usullari bilan birga



ishlatilganda, samaradorlik va hisoblash tezligi o'rtasida muvozanatni ta'minlash mumkin.

XULOSA

Mazkur maqolada shaxsni ovozi orqali identifikatsiyalash masalasi uchun MFCC, Δ , $\Delta\Delta$ va Mel-spektral belgilar asosida GMM-UBM modeli ishlab chiqildi va turli belgilar to'plamlarini kamaytirish usullari sinovdan o'tkazildi. Tajribalar shuni ko'rsatdiki:

1. Belgilarni kamaytirmasdan foydalanish o'rtacha natija berdi (67.4%).
2. PCA va ICA aniqlikni pasaytirdi.
3. GA usuli 70.9% aniqlikka erishdi.
4. LDA usuli eng yaxshi natijani berdi: aniqlik 77.5%, F1-makro 78.7%, Top-5 aniqlik esa 91.4%.

Natijalar shuni ko'rsatadiki, GMM-UBM + LDA kombinatsiyasi hisoblash quvvati cheklangan qurilmalarda shaxsni ovozi orqali identifikatsiyalash masalasi uchun eng maqbul yondashuvlardan biridir. Kelgusida ushbu tadqiqotni yanada rivojlantirish uchun boshqa optimallashtirish algoritmlarini (masalan, Particle Swarm Optimization, Ant Colony Optimization) qo'llash va real vaqt rejimida sinovdan o'tkazish rejalashtirilmoqda.

References:

1. Reynolds, D. A., Quatieri, T. F., & Dunn, R. B. (2000). Speaker verification using adapted Gaussian mixture models. *Digital Signal Processing*, 10(1-3), 19-41. <https://doi.org/10.1006/dspr.1999.0361>
2. Davis, S., & Mermelstein, P. (1980). Comparison of parametric representations for monosyllabic word recognition in continuously spoken sentences. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 28(4), 357-366. <https://doi.org/10.1109/TASSP.1980.1163420>
3. Young, S., Evermann, G., Gales, M., Kershaw, D., Liu, X., Moore, G., ... & Woodland, P. (2006). *The HTK Book (for HTK Version 3.4)*. Cambridge University Engineering Department.
4. Jolliffe, I. T., & Cadima, J. (2016). Principal component analysis: A review and recent developments. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 374(2065), 20150202. <https://doi.org/10.1098/rsta.2015.0202>
5. Hyvärinen, A., & Oja, E. (2000). Independent component analysis: Algorithms and applications. *Neural Networks*, 13(4-5), 411-430. [https://doi.org/10.1016/S0893-6080\(00\)00026-5](https://doi.org/10.1016/S0893-6080(00)00026-5)
6. Fisher, R. A. (1936). The use of multiple measurements in taxonomic problems. *Annals of Eugenics*, 7(2), 179-188. <https://doi.org/10.1111/j.1469-1809.1936.tb02137.x>
7. Goldberg, D. E. (1989). *Genetic Algorithms in Search, Optimization and Machine Learning*. Addison-Wesley.
8. Kinnunen, T., & Li, H. (2010). An overview of text-independent speaker recognition: From features to supervectors. *Speech Communication*, 52(1), 12-40. <https://doi.org/10.1016/j.specom.2009.08.009>



9. Rabiner, L., & Juang, B. H. (1993). Fundamentals of Speech Recognition. Prentice-Hall.
10. Bishop, C. M. (2006). Pattern Recognition and Machine Learning. Springer.