

Phonexia Voice Inspector

Powered by Phonexia Deep Embeddings™

“With suitable speech materials, speaker comparisons are carried out by means of modern and validated voice biometric methods. For that purpose, stimmenvergleich.com uses cutting-edge language-, text- and channel-independent Phonexia speaker identification software.

Dr. Stefan Gfroerer

Former head of the Speaker Identification and Audio Analysis Department at Bundeskriminalamt

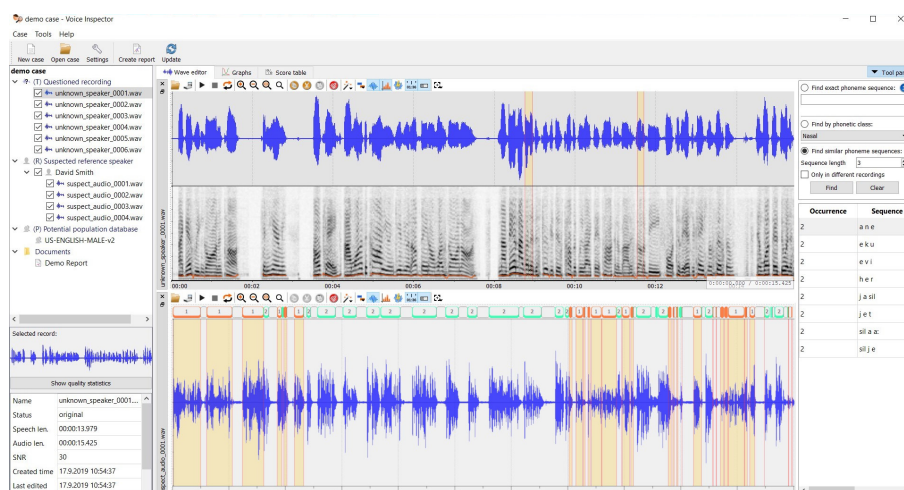
Phonexia Voice Inspector (VIN) provides police forces and forensic experts with a highly accurate speaker identification tool to support criminal investigations. It leverages the AI-powered Phonexia Deep Embeddings™ voice biometrics engine which is, based on an independent forensic study (forensic_eval_01) performed by Bundeskriminalamt, the most accurate automatic speaker identification technology available on the market.

Selected Features

- **1:1 speaker comparison** in accordance with ENFSI guidelines
- **1:N speaker identification** for more complex cases
- **A diarization tool** to make working with audio recordings containing multiple speakers easier
- **A phoneme recognizer** for the searching and visualization of the same phoneme sequences across audio files
- **An evaluation tool** for the measurement of accuracy in a user's data sets
- **A waveform editor** with tools such as a spectrum panel, voice activity detection and more

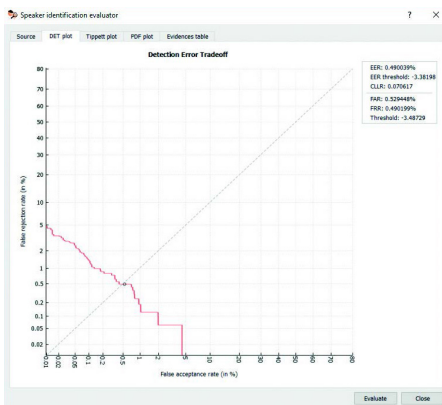
Turning Voice into Knowledge

phonexia.com



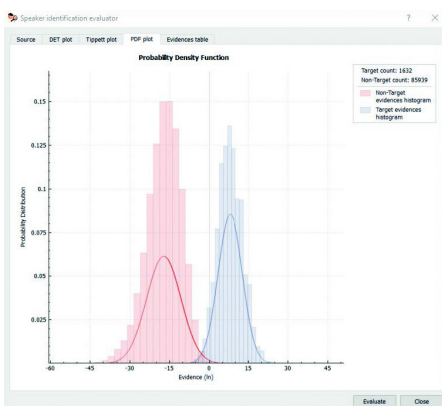
Automatic Forensic Voice Comparison

The *forensic_eval_01* study's details
<https://phonexia.com/en/forensic-report>



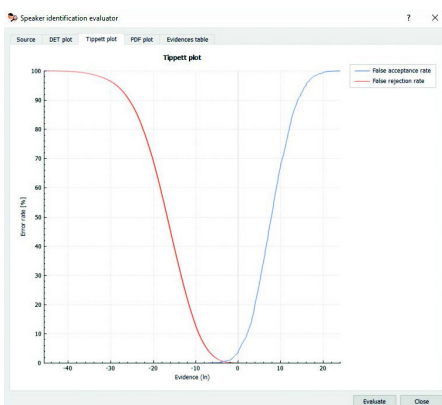
Technology

- **Deep Embeddings™** – uses deep neural networks to generate highly representative voiceprints
- **Phonexia Voice Inspector** is independent of language, accent, text and channel
- Applies state-of-the-art **channel compensation techniques**, verified by NIST evaluation
- **Compatible** with the widest range of audio sources possible: GSM/CDMA, 3G, VoIP, landlines, etc.



Input

- **Input format for processing:**
WAV or RAW (8 or 16-bit linear coding), A-law or Mu-law, PCM, 8 kHz+ sampling
- **Minimum speech signal for enrollment:**
Recommended 20+ seconds
- **Minimum speech signal for identification:**
Recommended 7+ seconds



Output

- **Scoring** to a likelihood ratio (LR), log-likelihood ratio (LLR) and verbal presentation of results
- **Graphic presentation** of the likelihood ratio (LR)
- **Detailed report output** (expert opinion template automatically generated) for presentation of results (to a court or an investigation team)

Turning Voice
into Knowledge

phonexia.com

Chaloupkova 3002/1a, 612 00 Brno
Czech Republic, European Union
info@phonexia.com
+420 511 205 265

