

VoiceRadar: Voice Deepfake Detection using Micro-Frequency and Compositional Analysis

Kavita Kumari, Maryam Abbasihafshejani,
Alessandro Pegoraro, Phillip Rieger, Kamyar Arshi,
Murtuza Jadliwala, Ahmad-Reza Sadeghi

NDSS 2025



TECHNISCHE
UNIVERSITÄT
DARMSTADT



The University of Texas
at San Antonio™



Audio Deepfakes

Fraud



Business

Unusual CEO Fraud via Deepfake Audio
Steals US\$243,000 From UK Company

September 05, 2019

Audio Deepfakes

Fraud



Business

Unusual CEO Fraud via Deepfake Audio Steals US\$243,000 From UK Company

September 05, 2019

Disinformation



Politics

SCOTUS

Congress

Facts First

2024 Elections



INVESTIGATES

A fake recording of a candidate saying he'd rigged the election went viral. Experts say it's only the beginning

Audio Deepfakes

Defamation



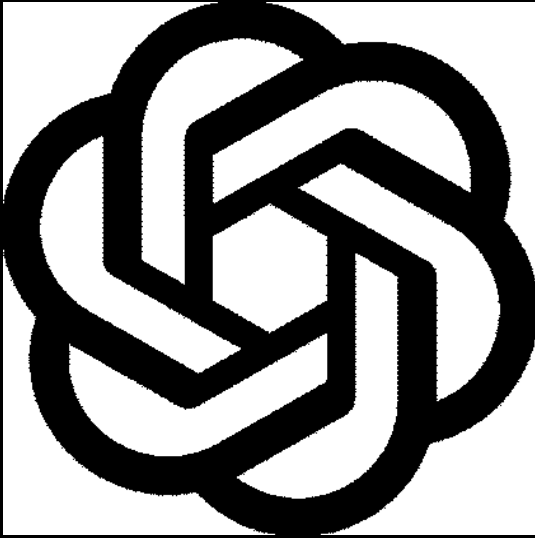
U.S. NEWS

Deepfake of principal's voice is the latest case of AI being used for harm



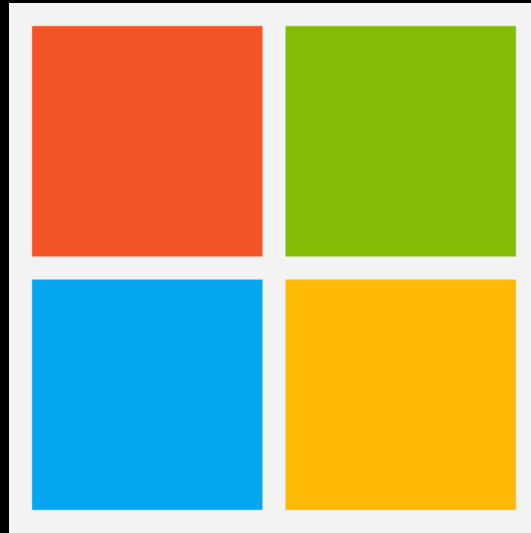
The Audio Generation Arms Race

OpenAI



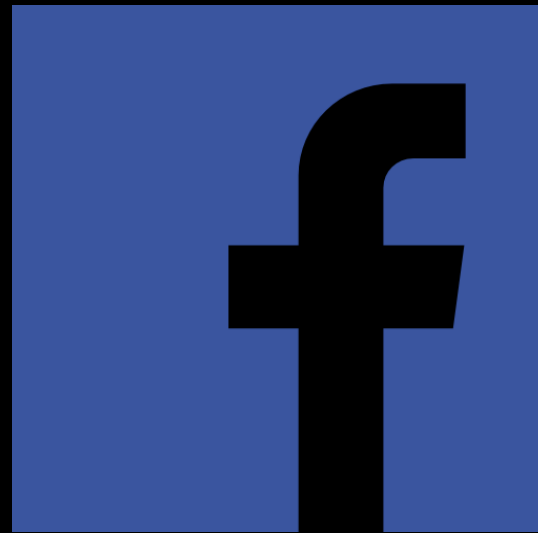
Voice Engine
(Voice Cloning with
15 sec of samples)
[Open AI Blog]

Microsoft



VALLE-X
(Text-to-speech voice generation
using Language Modelling)
[Meng et al. arXiv 2024]

Facebook/Meta



Singing VC
(Singing Voice Conversion)
[Nachmani et al. Interspeech 2019]

Google



AudioLM
(Audio Generation using
Language Modelling)
[Borsos et al. arXiv 2022]

General Principle of Deepfake



Original Audio Data

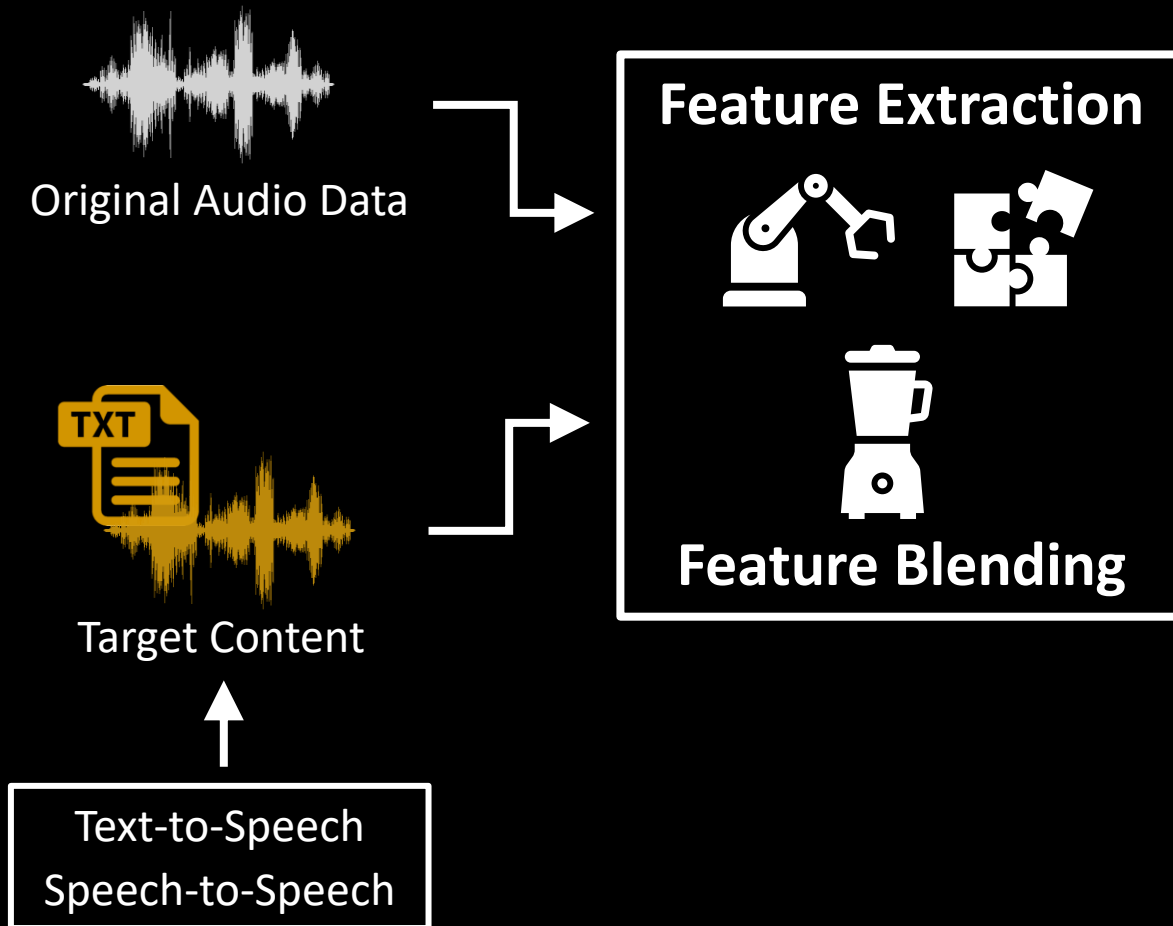


Target Content

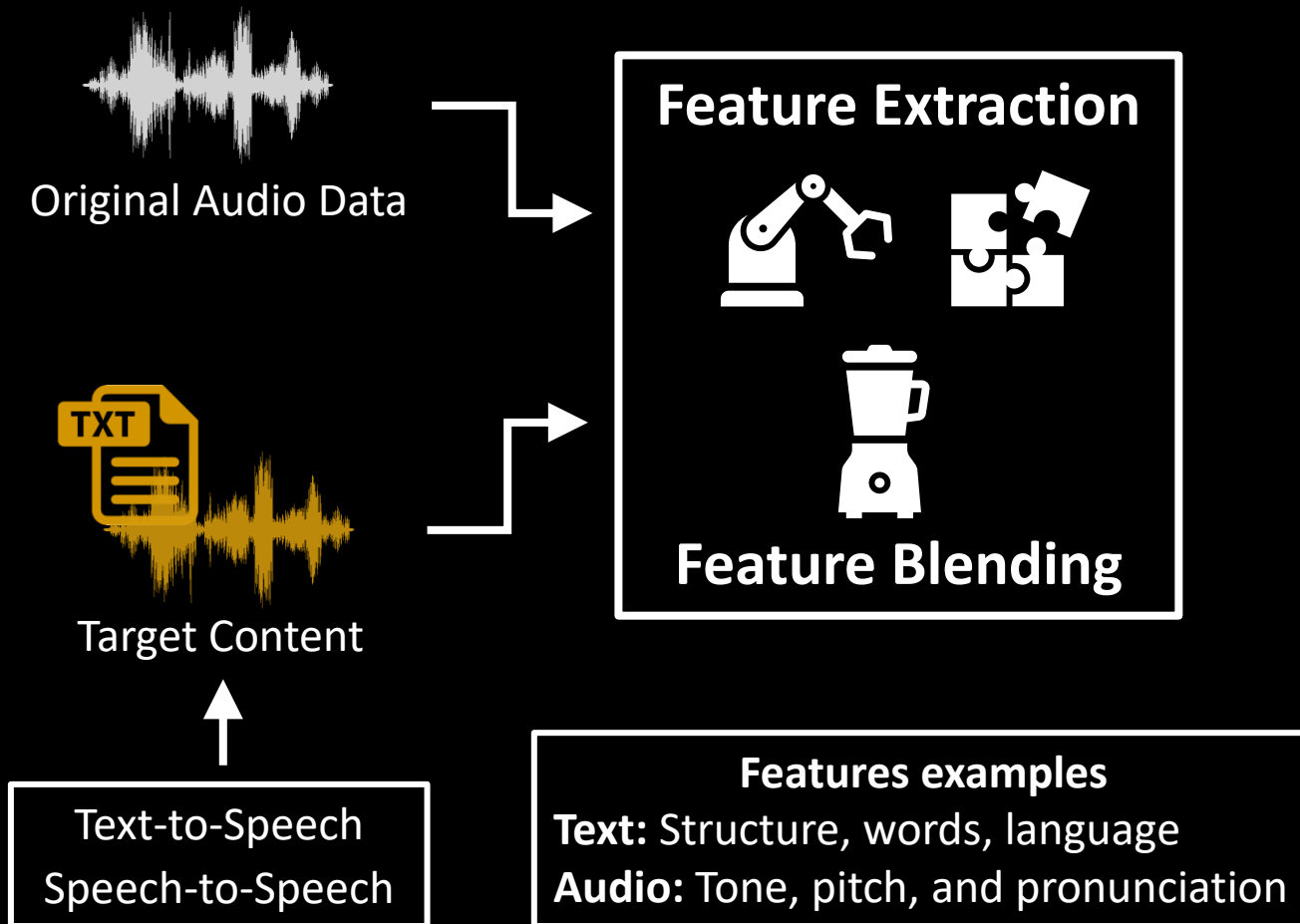


Text-to-Speech
Speech-to-Speech

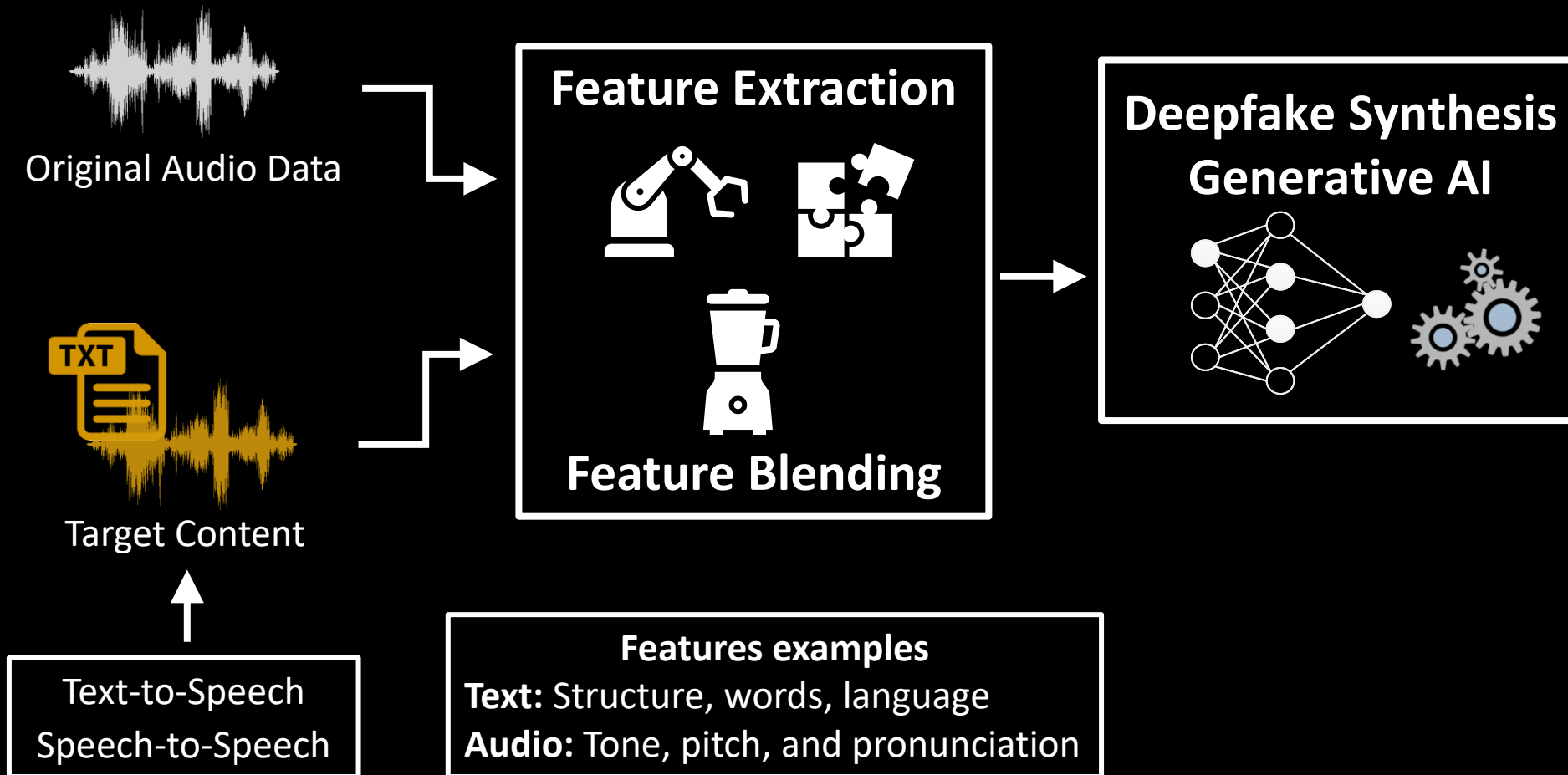
General Principle of Deepfake



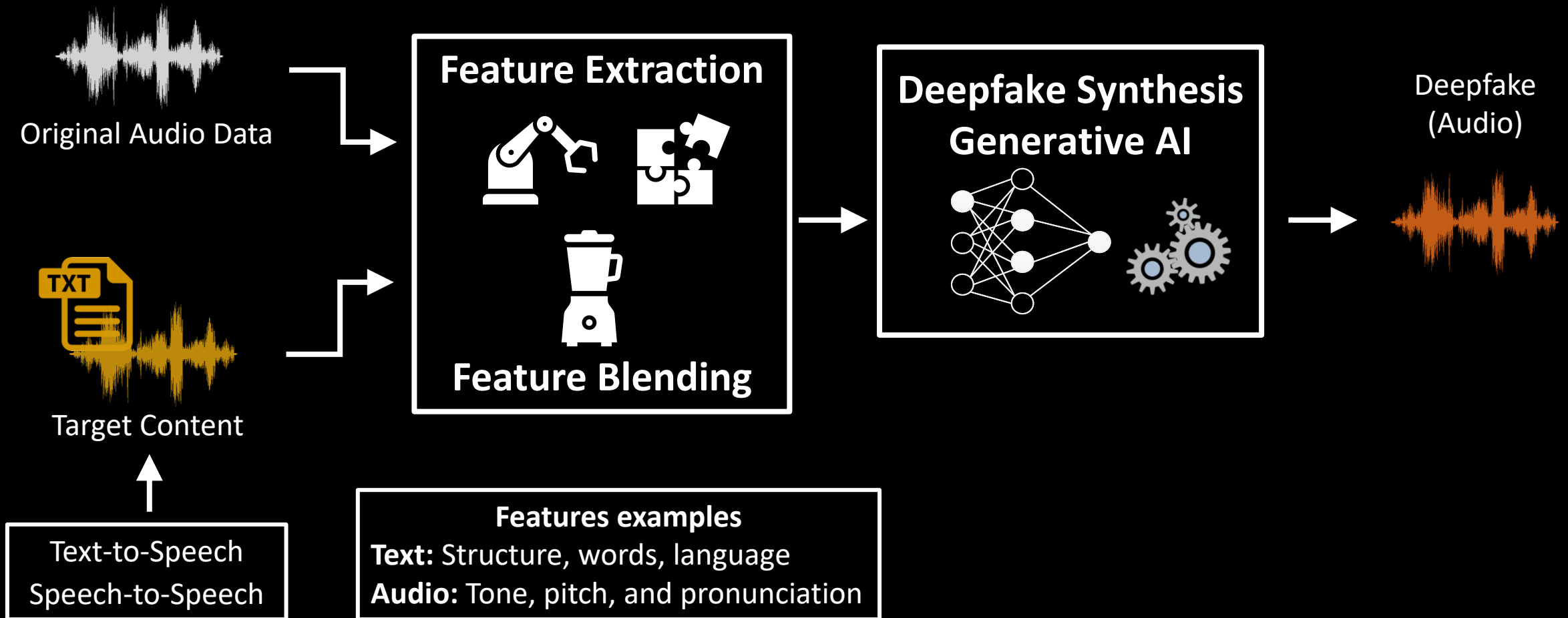
General Principle of Deepfake



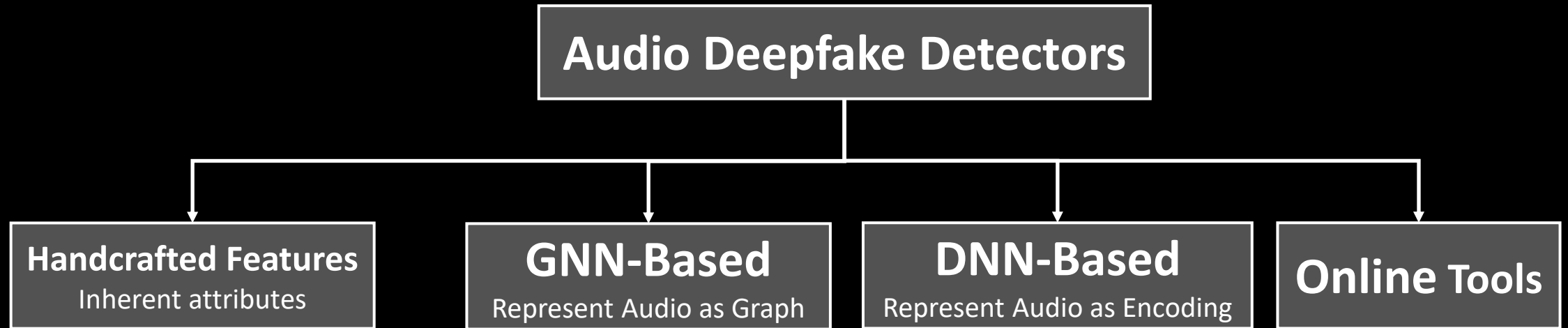
General Principle of Deepfake



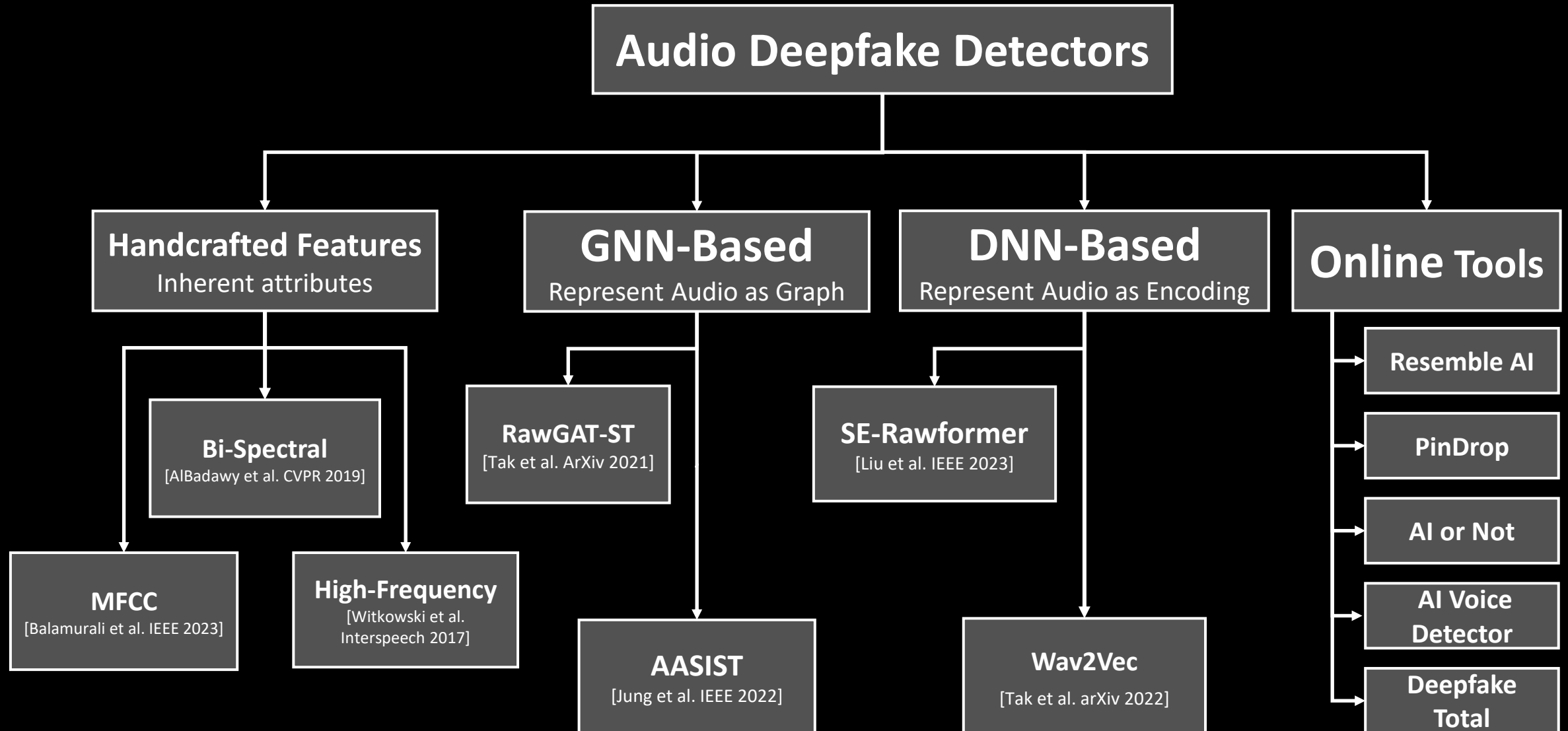
General Principle of Deepfake



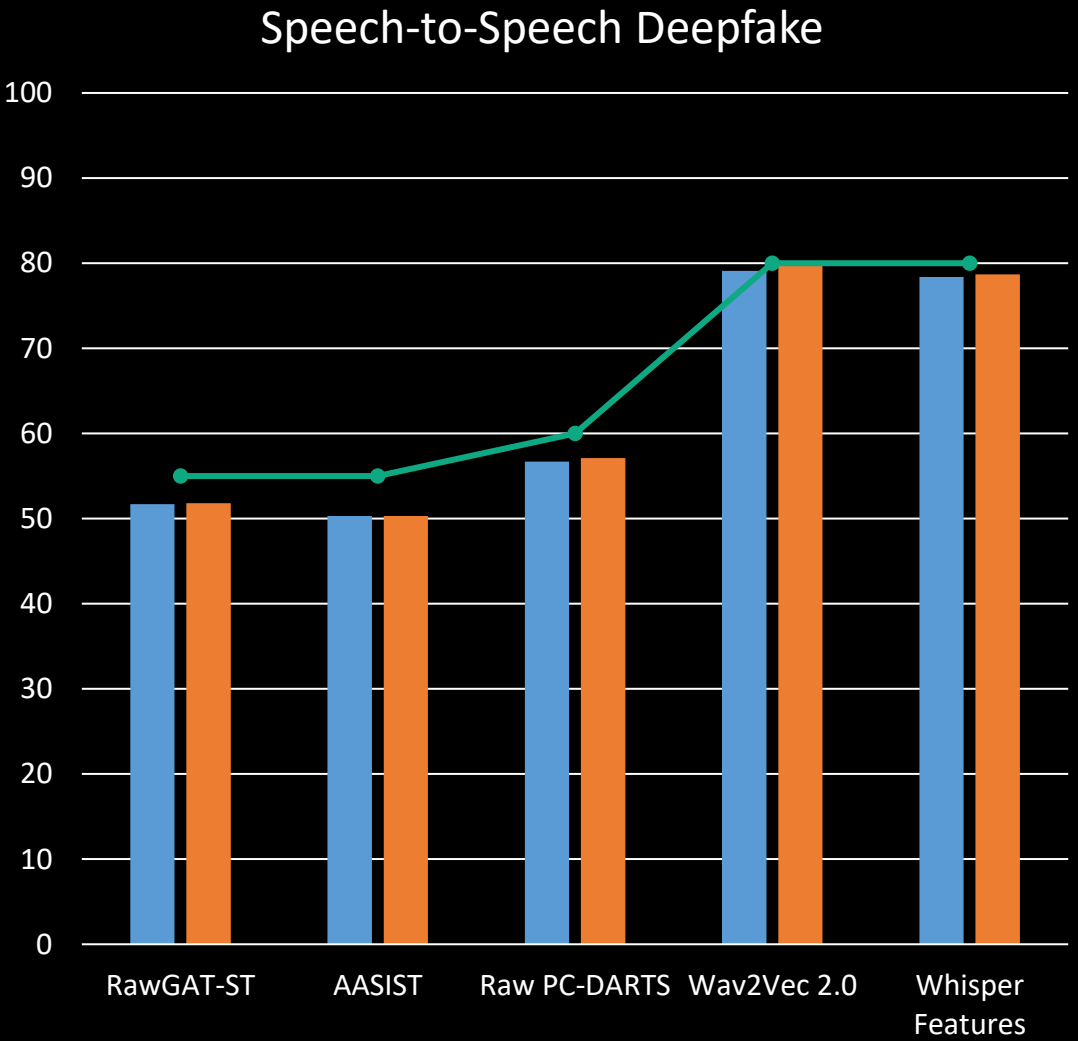
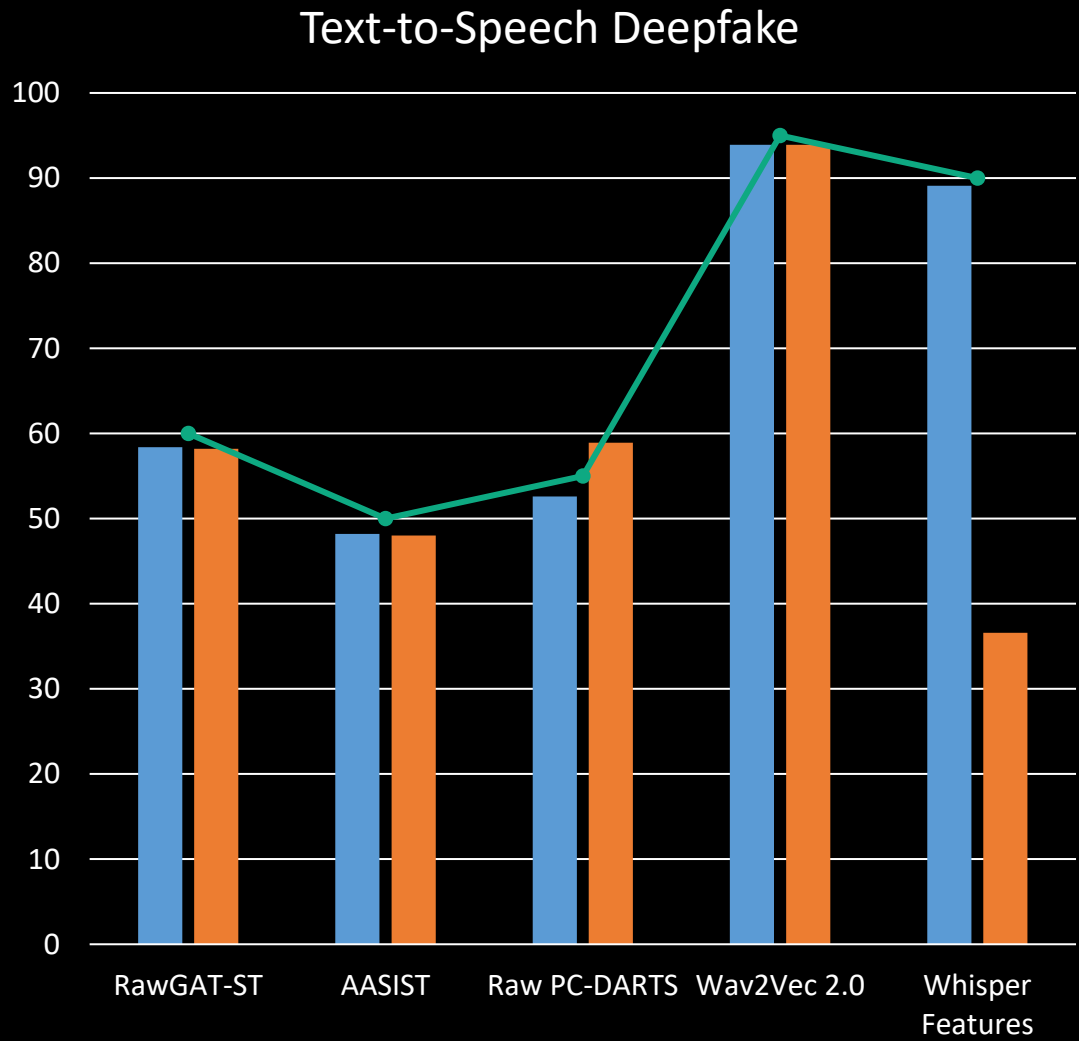
Categorization of Audio Deepfake Detectors



Categorization of Audio Deepfake Detectors



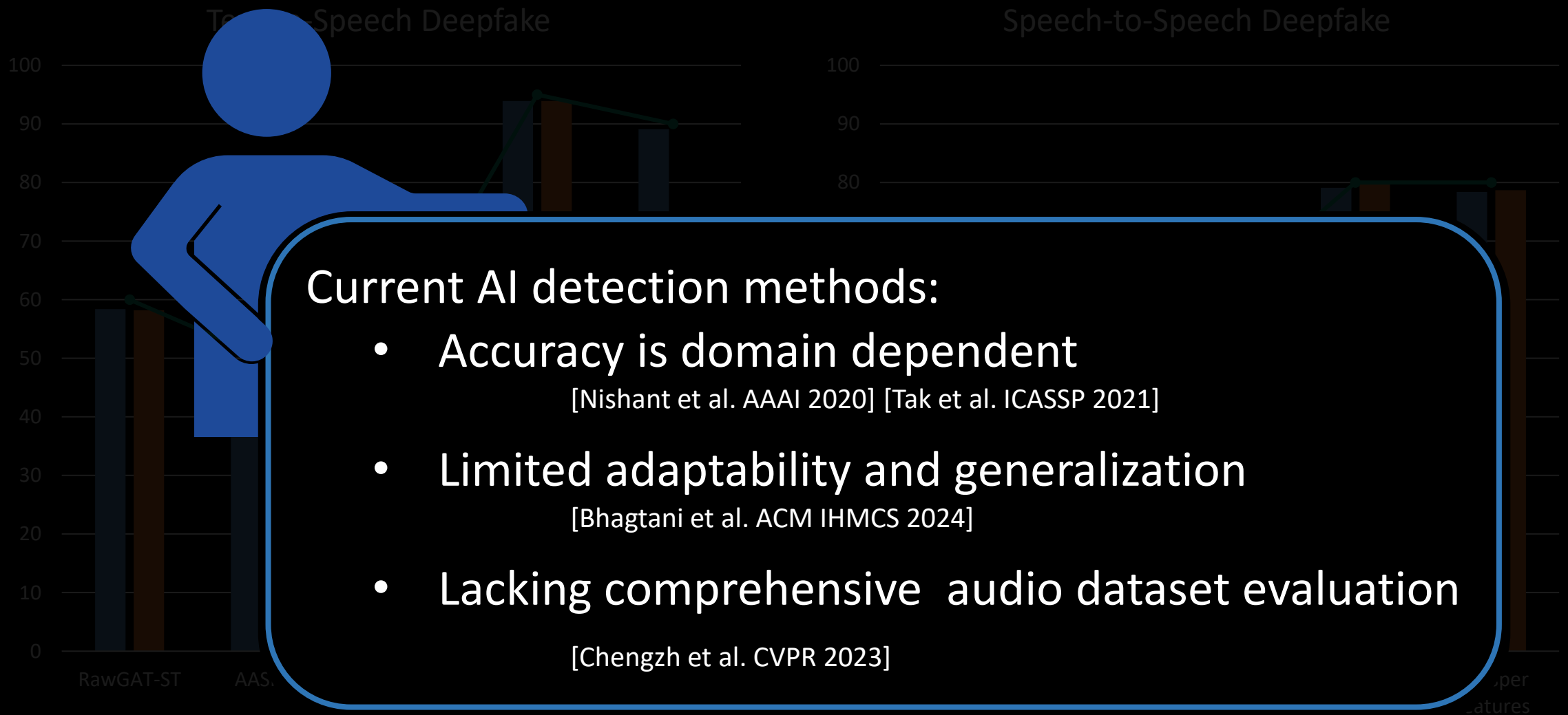
Evaluation of Deepfake Detectors



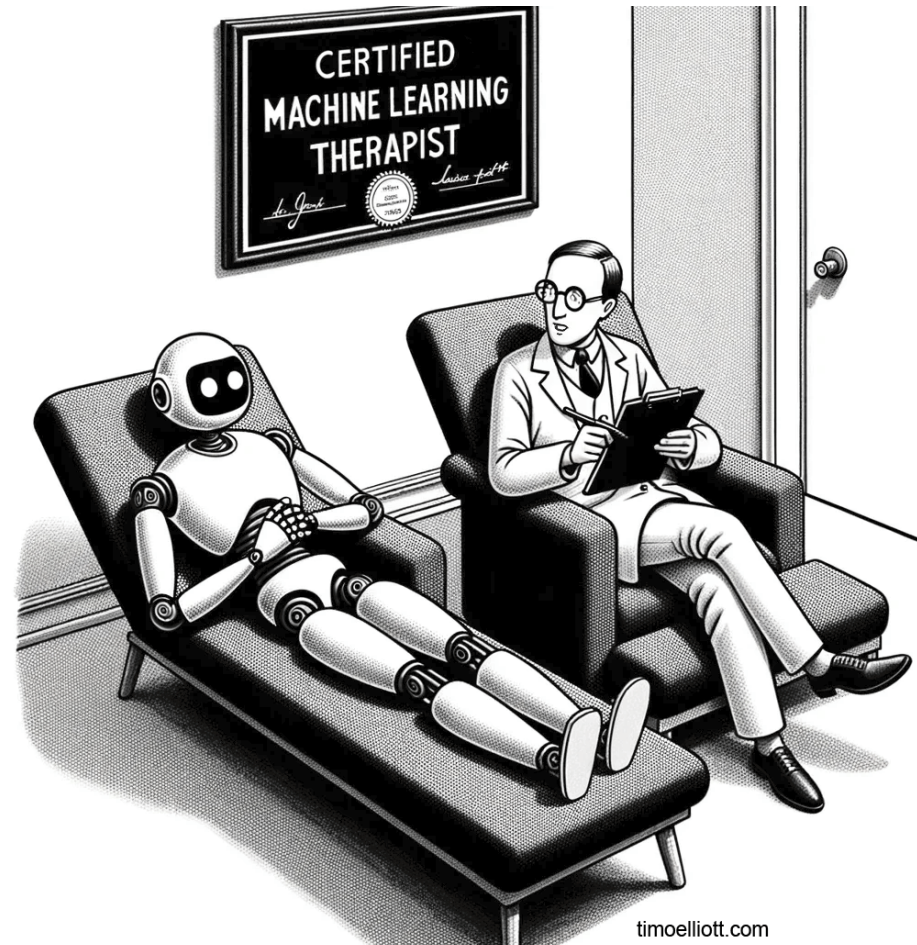
$$TPR = \frac{TP}{TP + FN}$$

$$TNR = \frac{TN}{TN + FP}$$

Evaluation of Deepfake Detectors



Human vs. Machine

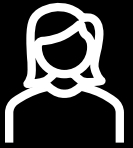


timoelliott.com

“So, when did you first realize humans might have feelings?”

Intuition and Hypothesis

Human interactions

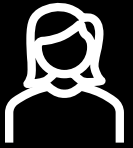


- Affected by personal and environmental factors dependent on individuals involved



Intuition and Hypothesis

Human interactions

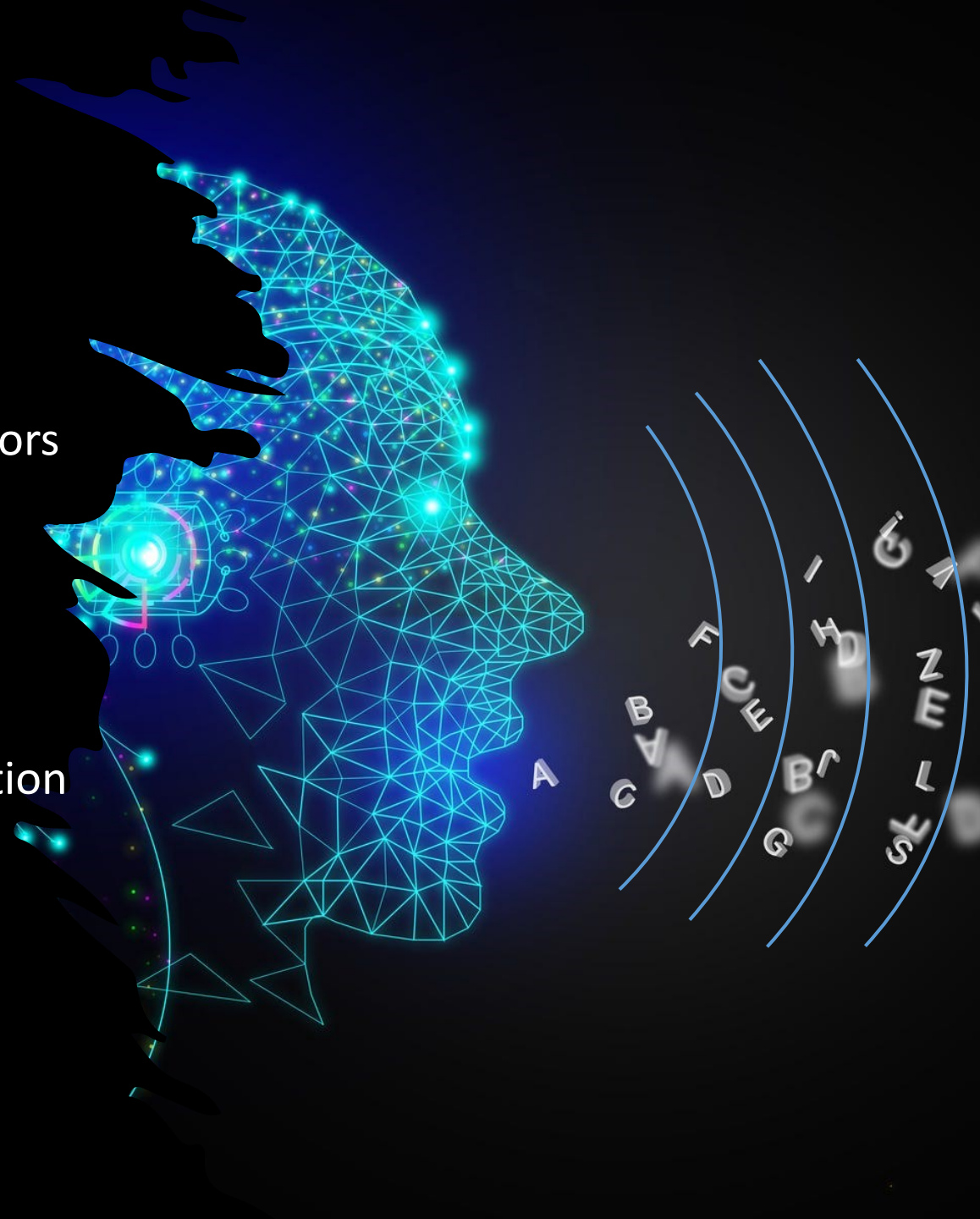


- Affected by personal and environmental factors dependent on individuals involved

Machines interactions



- Follow predefined rules and pattern recognition
- Less dynamic than human interactions





VoiceRadar

VoiceRadar: Motivation

Approximates the
physical models

VoiceRadar: Motivation

Approximates the
physical models



To simulate the
audio wave
propagation

VoiceRadar: Motivation

Approximates the
physical models

Extracts the subtle
micro-frequencies



To simulate the
audio wave
propagation

VoiceRadar: Motivation

Approximates the
physical models



To simulate the
audio wave
propagation

Extracts the subtle
micro-frequencies



To differentiate
synthetic audio

VoiceRadar: Motivation

Approximates the
physical models



To simulate the
audio wave
propagation

Extracts the subtle
micro-frequencies



To differentiate
synthetic audio

Incorporates
compositional bias

VoiceRadar: Motivation

Approximates the
physical models



To simulate the
audio wave
propagation

Extracts the subtle
micro-frequencies



To differentiate
synthetic audio

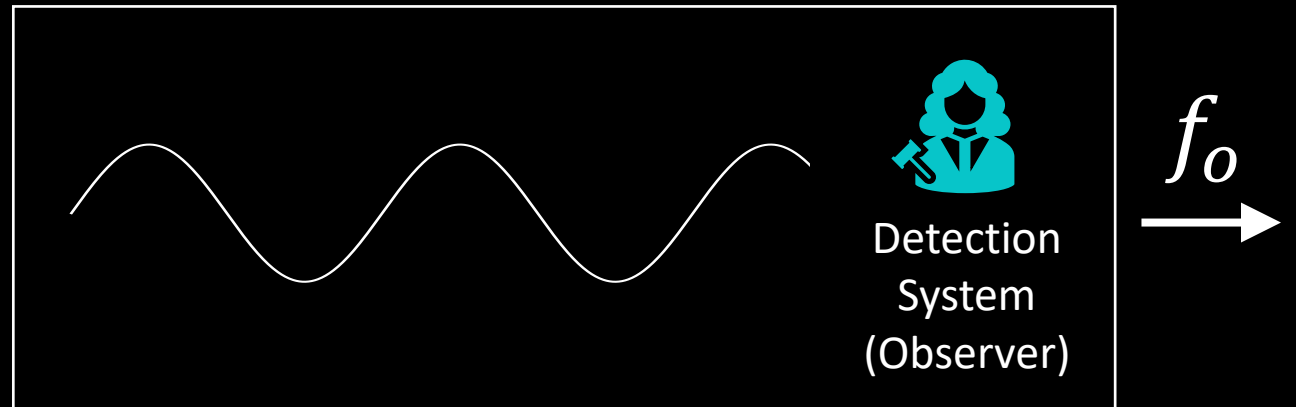
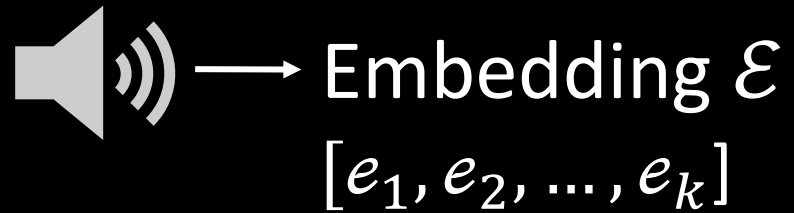
Incorporates
compositional bias



To penalize mismatch
of energies by the
model

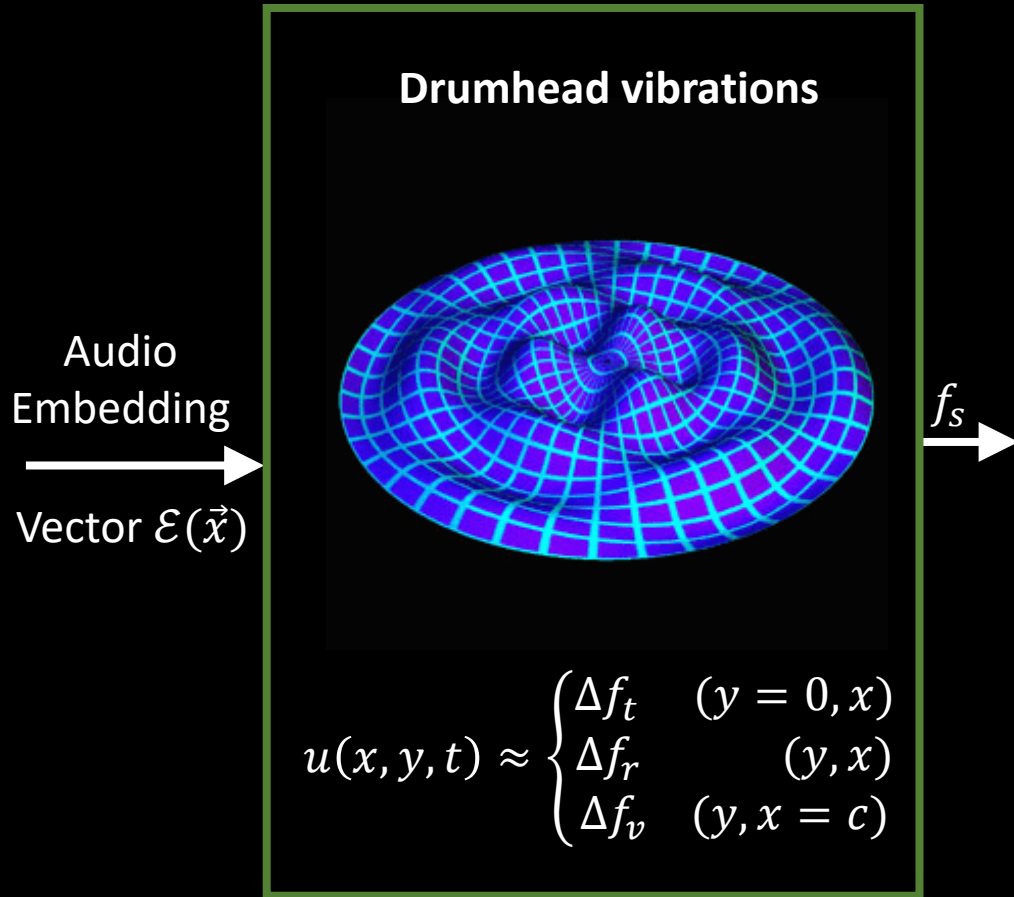
VoiceRadar: Micro-Frequency Analysis

[Observer Frequency]

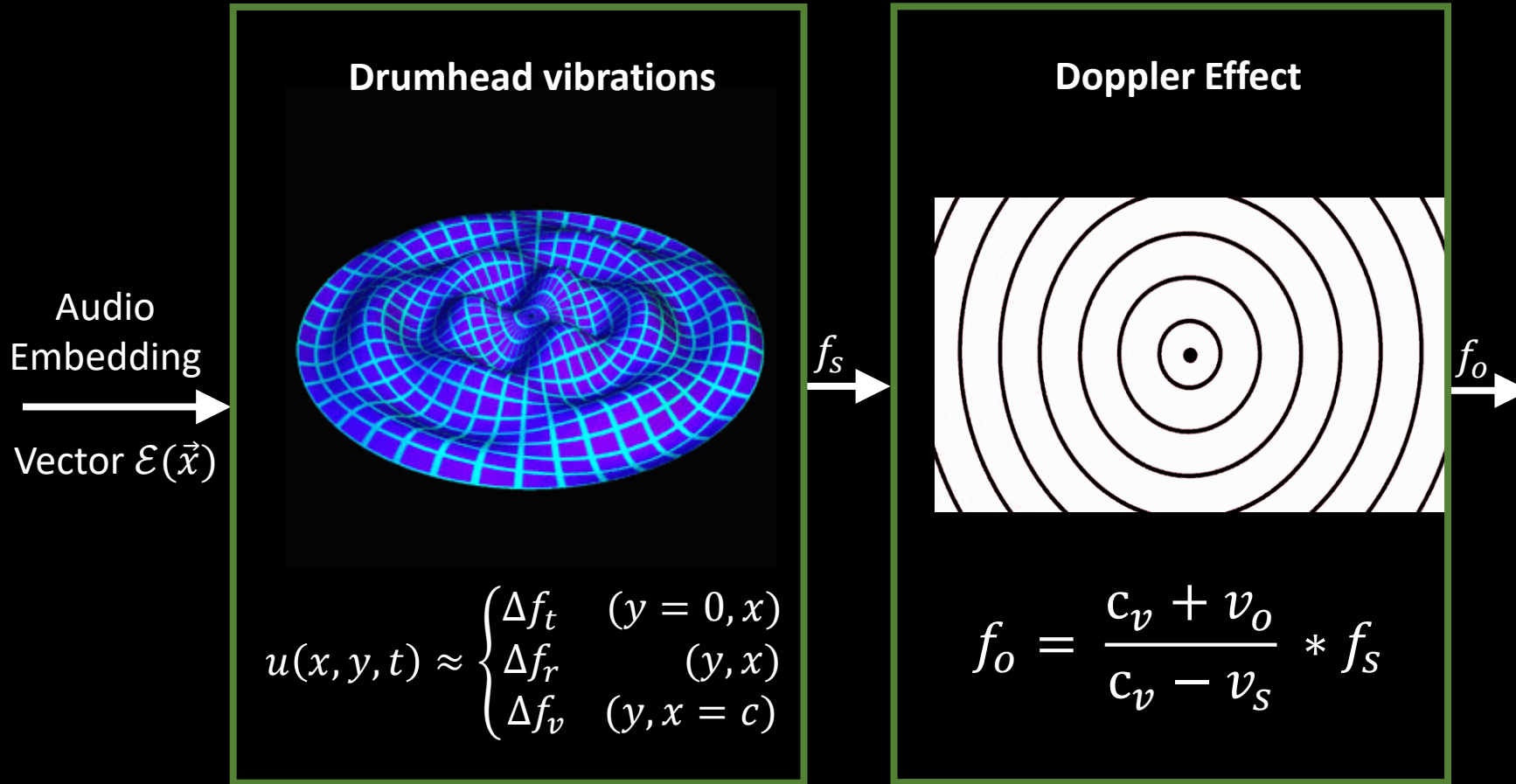


Physical models approximation

VoiceRadar: Core Components



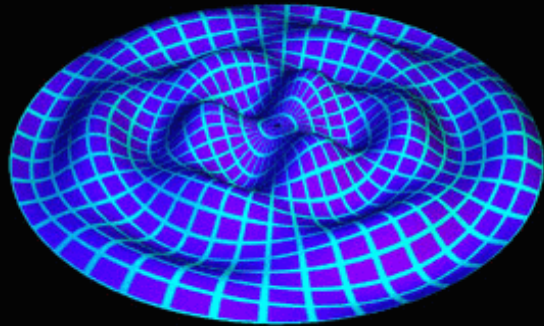
VoiceRadar: Core Components



VoiceRadar: Core Components

Audio
Embedding
→
Vector $\mathcal{E}(\vec{x})$

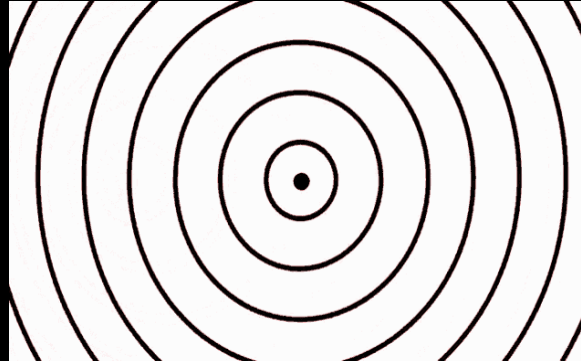
Drumhead vibrations



$$u(x, y, t) \approx \begin{cases} \Delta f_t & (y = 0, x) \\ \Delta f_r & (y, x) \\ \Delta f_v & (y, x = c) \end{cases}$$

f_s

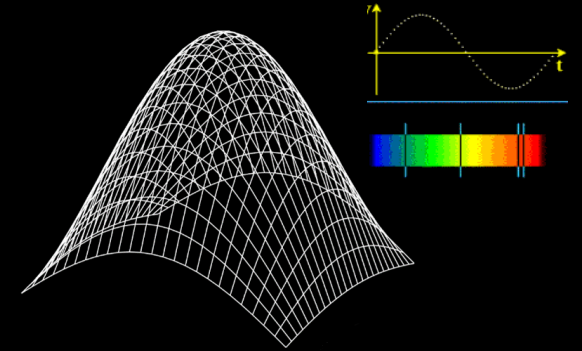
Doppler Effect



f_o

$$f_o = \frac{c_v + v_o}{c_v - v_s} * f_s$$

Model Training with Physics Prior



$$L = \text{Binay_cross_entropy} + f_o$$
$$W : \text{argmin}_W L(W, f_o)$$

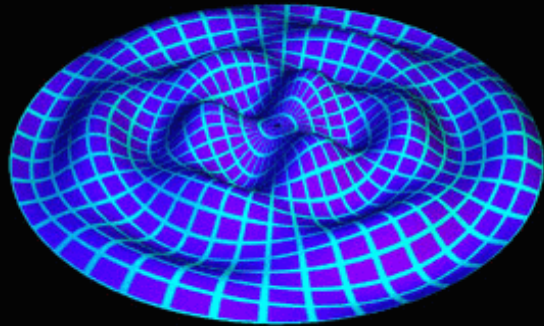


VoiceRadar: Core Components

Embedding $\mathcal{E}(\vec{x}) := [e_1, e_2, \dots, e_k]$

Audio
Embedding
→
Vector $\mathcal{E}(\vec{x})$

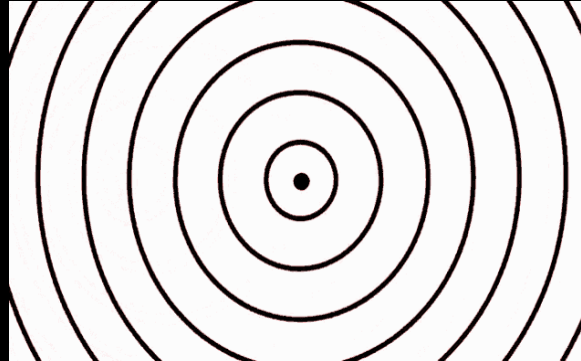
Drumhead vibrations



$$u(x, y, t) \approx \begin{cases} \Delta f_t & (y = 0, x) \\ \Delta f_r & (y, x) \\ \Delta f_v & (y, x = c) \end{cases}$$

f_s

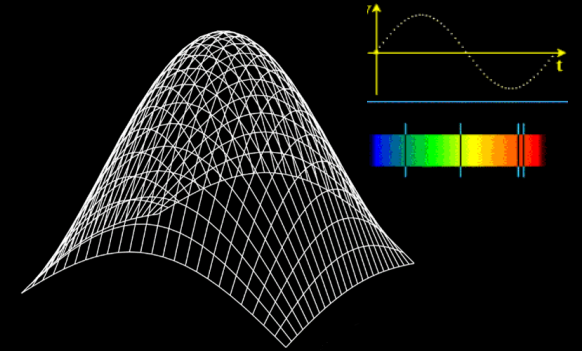
Doppler Effect



f_o

$$f_o = \frac{c_v + v_o}{c_v - v_s} * f_s$$

Model Training
with Physics Prior



$$L = \text{Binay_cross_entropy} + f_o$$
$$W : \operatorname{argmin}_W L(W, f_o)$$

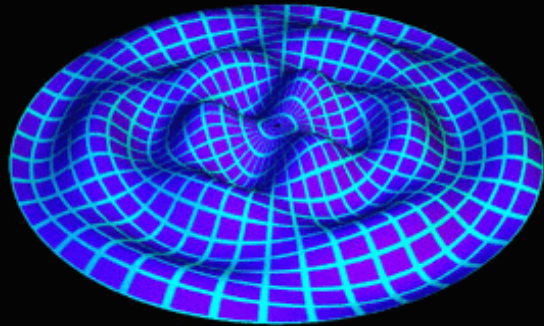


VoiceRadar: Core Components

Embedding $\mathcal{E}(\vec{x}) := [e_1, e_2, \dots, e_k]$

Audio
Embedding
Vector $\mathcal{E}(\vec{x})$

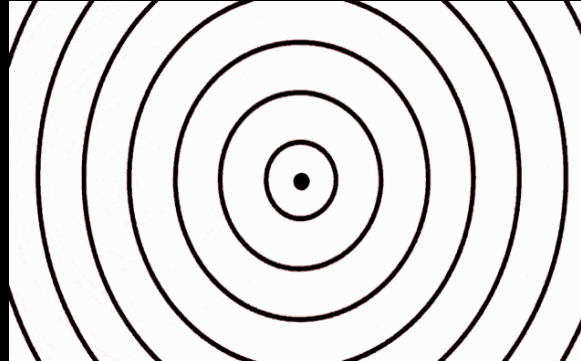
Drumhead vibrations



$$u(x, y, t) \approx \begin{cases} \Delta f_t & (y = 0, x) \\ \Delta f_r & (y, x) \\ \Delta f_v & (y, x = c) \end{cases}$$

f_s

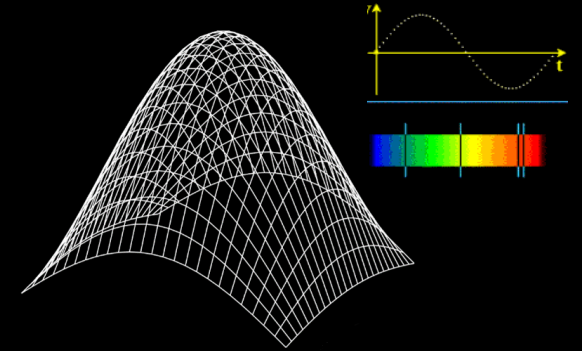
Doppler Effect



f_o

$$f_o = \frac{c_v + v_o}{c_v - v_s} * f_s$$

Model Training
with Physics Prior



$$L = \text{Binay_cross_entropy} + f_o$$
$$W : \text{argmin}_W L(W, f_o)$$

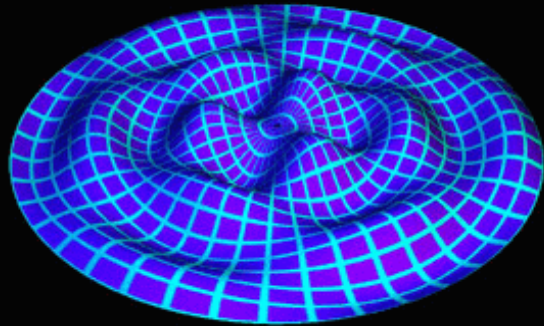


VoiceRadar: Core Components

Embedding $\mathcal{E}(\vec{x}) := [e_1, e_2, \dots, e_k]$

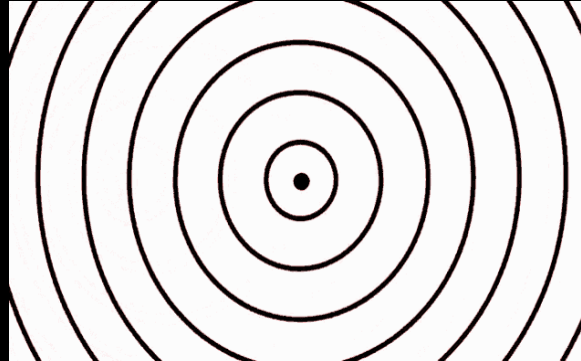
Audio
Embedding
Vector $\mathcal{E}(\vec{x})$

Drumhead vibrations



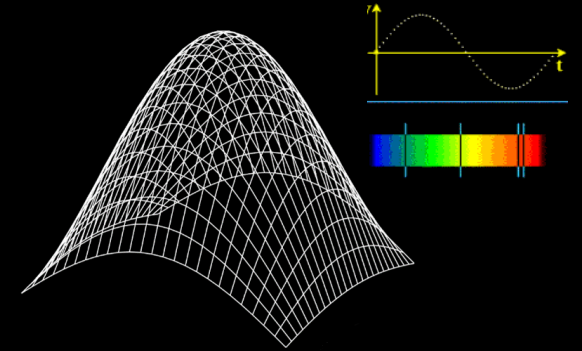
$$u(x, y, t) \approx \begin{cases} \Delta f_t & (y = 0, x) \\ \Delta f_r & (y, x) \\ \Delta f_v & (y, x = c) \end{cases}$$

Doppler Effect

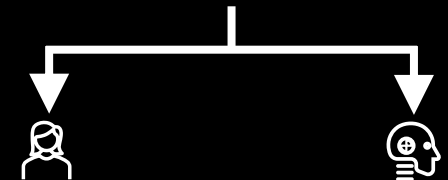


$$f_o = \frac{c_v + v_o}{c_v - v_s} * f_s$$

Model Training
with Physics Prior

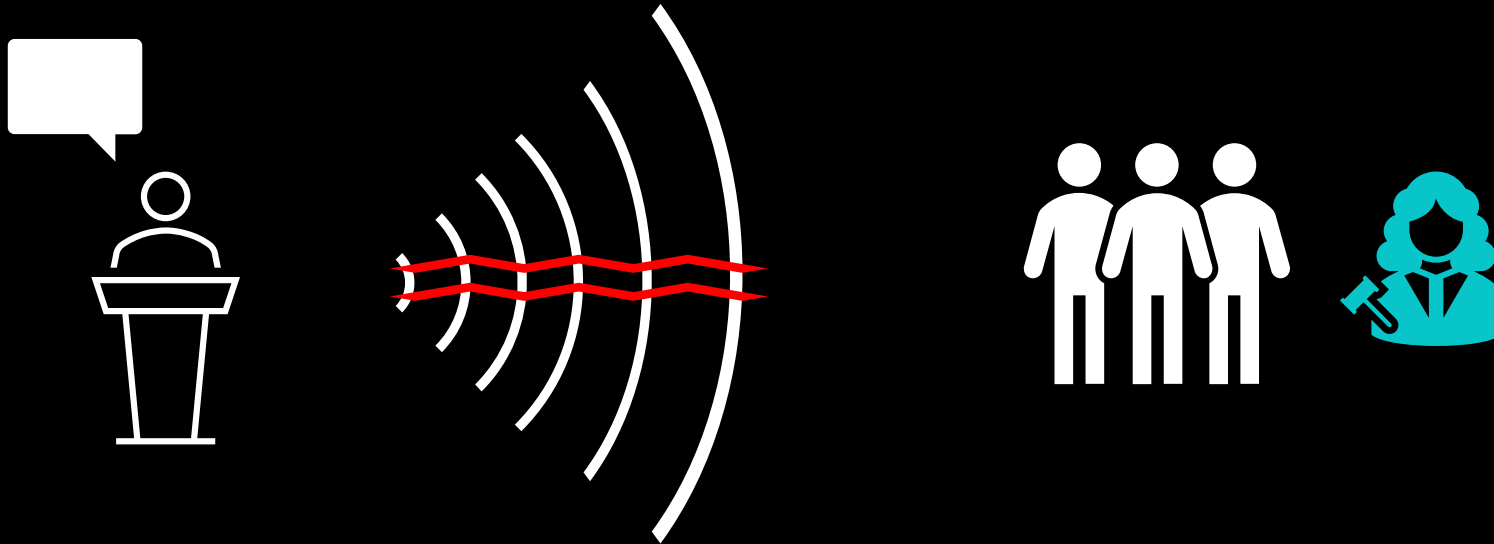


$$L = \text{Binay_cross_entropy} + f_o$$
$$W : \text{argmin}_W L(W, f_o)$$



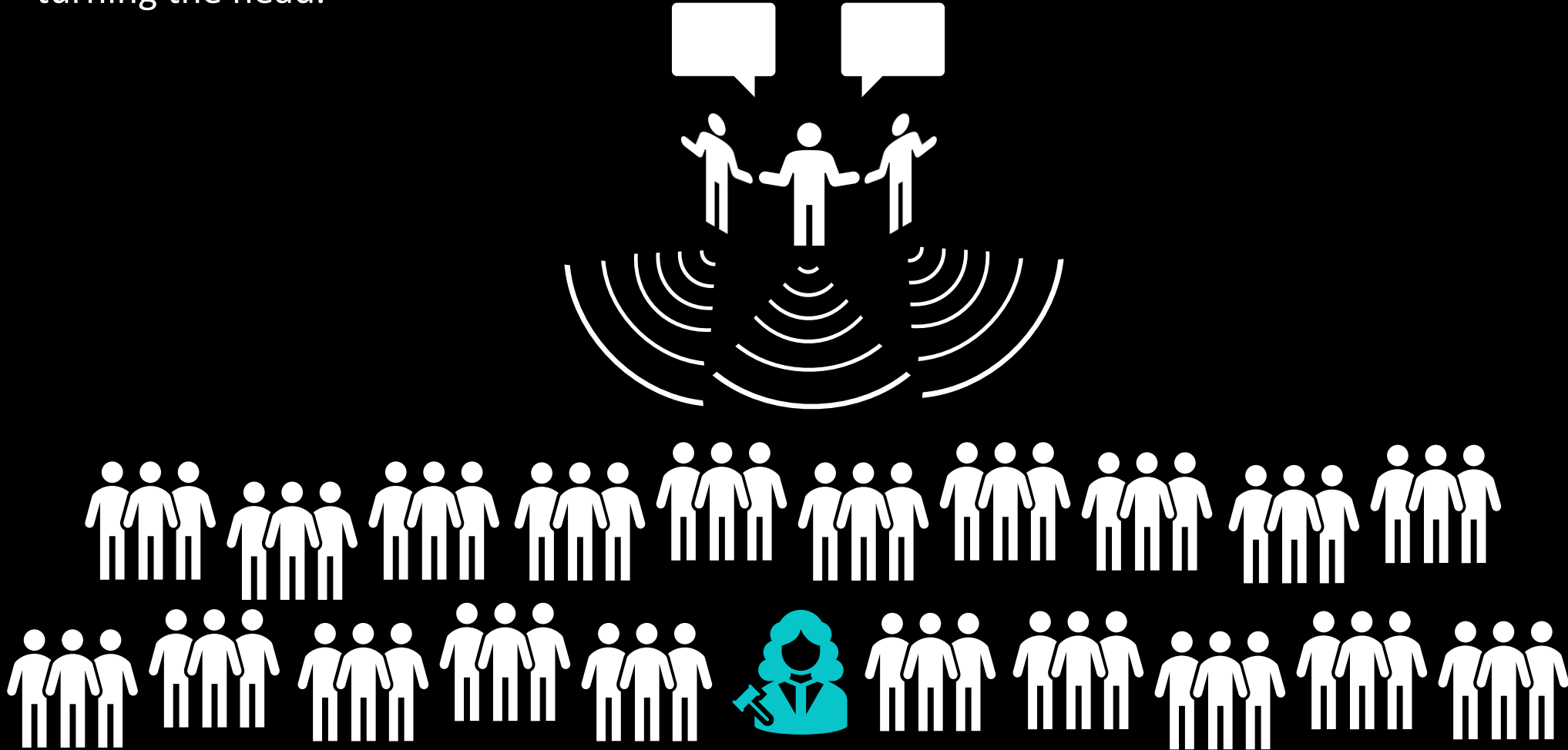
VoiceRadar: Translational Frequency

- Wave propagation of speech in straight lines from speaker to observer.



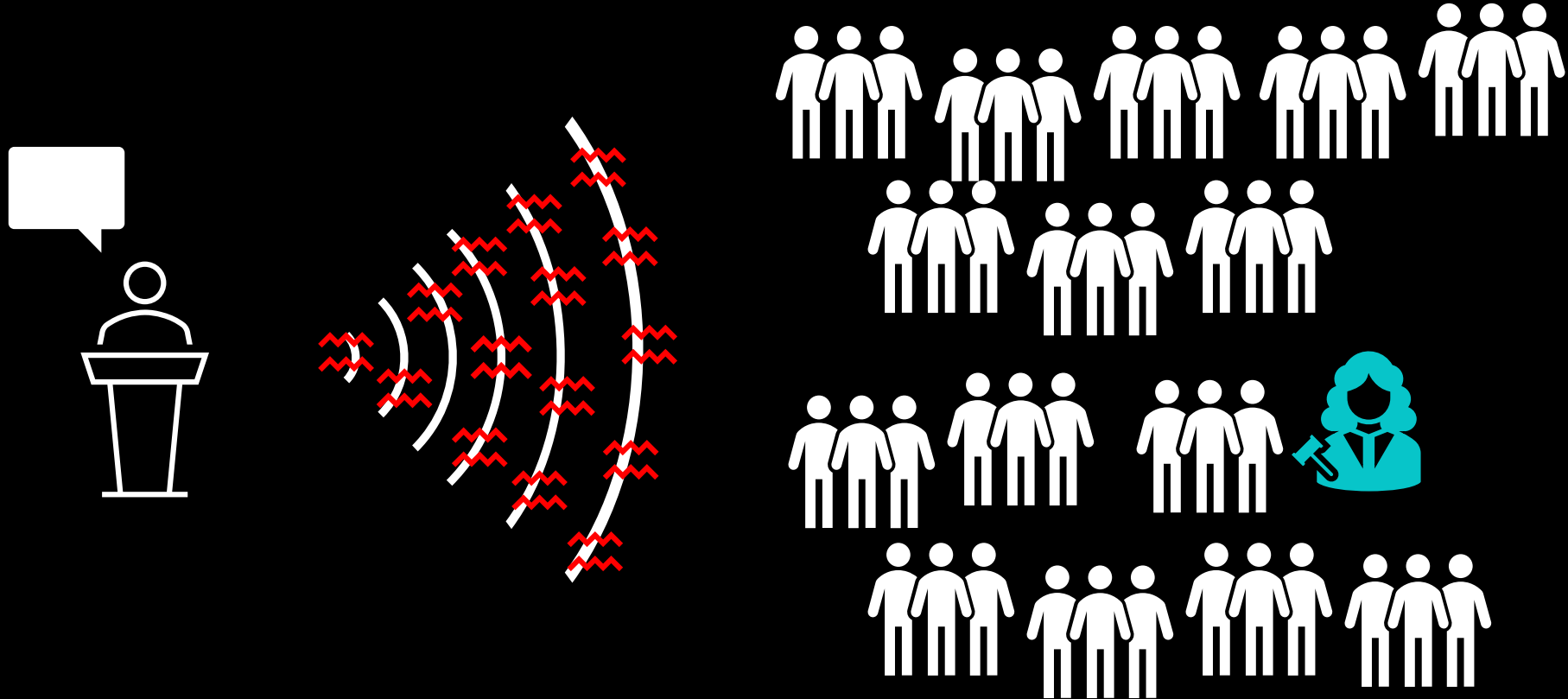
VoiceRadar: Rotational Frequency

- Voice projection changes while speaking causing a rotational frequency shift, like when turning the head.



VoiceRadar: Vibrational Frequency

- Speech has subtle vibrational variations
- Stemming from individual and environmental factors, causing tremors.

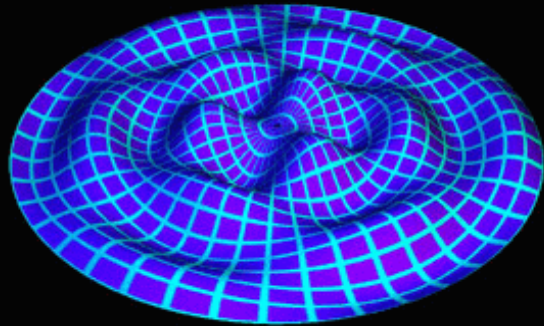


VoiceRadar: Core Components

Embedding $\mathcal{E}(\vec{x}) := [e_1, e_2, \dots, e_k]$

Audio
Embedding
Vector $\mathcal{E}(\vec{x})$

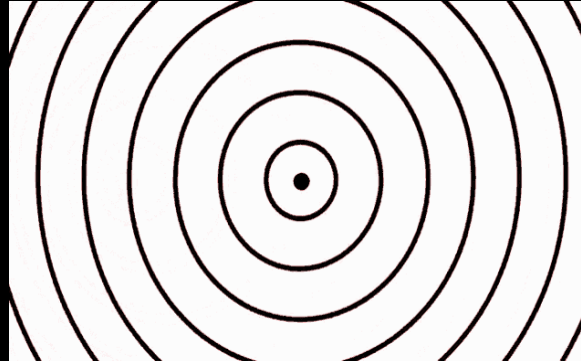
Drumhead vibrations



$$u(x, y, t) \approx \begin{cases} \Delta f_t & (y = 0, x) \\ \Delta f_r & (y, x) \\ \Delta f_v & (y, x = c) \end{cases}$$

f_s

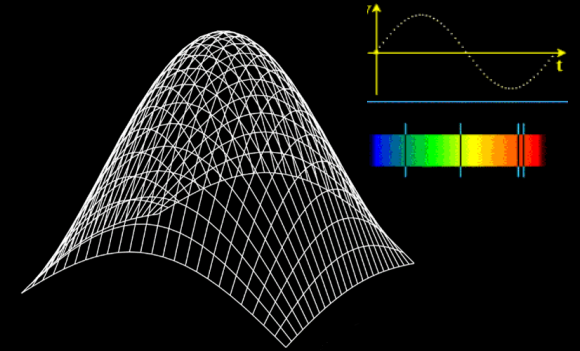
Doppler Effect



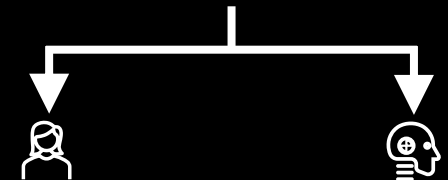
f_o

$$f_o = \frac{c_v + v_o}{c_v - v_s} * f_s$$

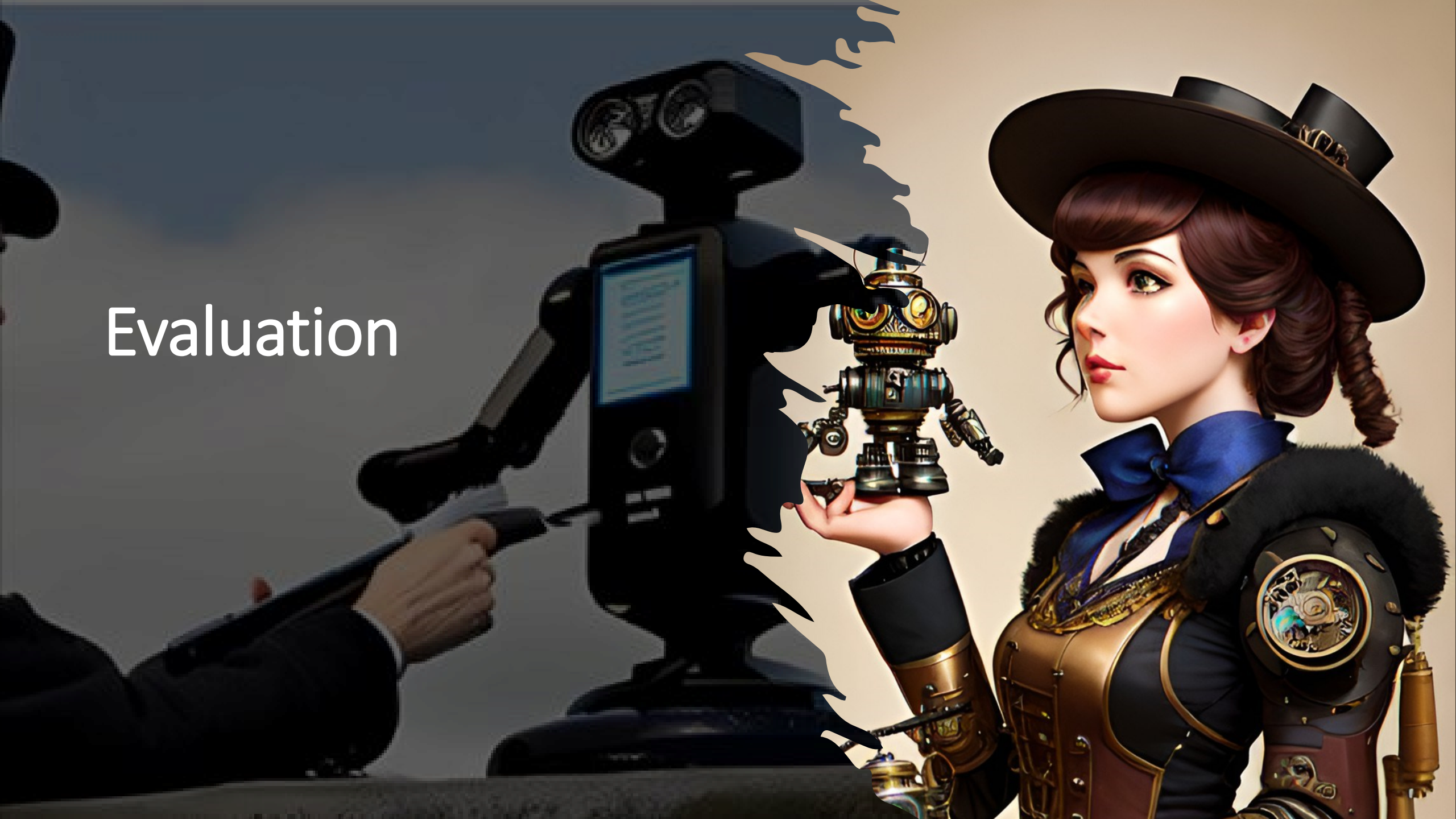
Model Training
with Physics Prior



$$L = \text{Binary_cross_entropy} + f_o$$
$$W : \text{argmin}_W L(W, f_o)$$



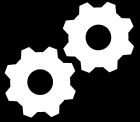
Evaluation



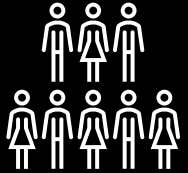
Text-to-Speech (TTS) Dataset



Based on recordings from VCTK datasets



Used 8 most recent TTS approaches



40 Speakers

- Cover different distributions of: Gender, age, and regions



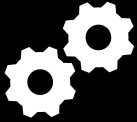
Generated 800 phrases for each speaker

Dataset	Records
VALL-E-X	32 000
Speech T5	32 000
Bark	32 000
Style TTS2	32 000
Jenny	32 040
Vits	32 040
XTTS	32 040
Tortoise	32 000
Total	256 120

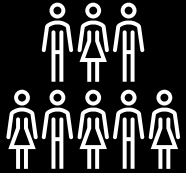
Speech-to-Speech (STS) Dataset



Based on recordings from VCTK datasets



Used 4 most recent STS approaches



40 Speakers

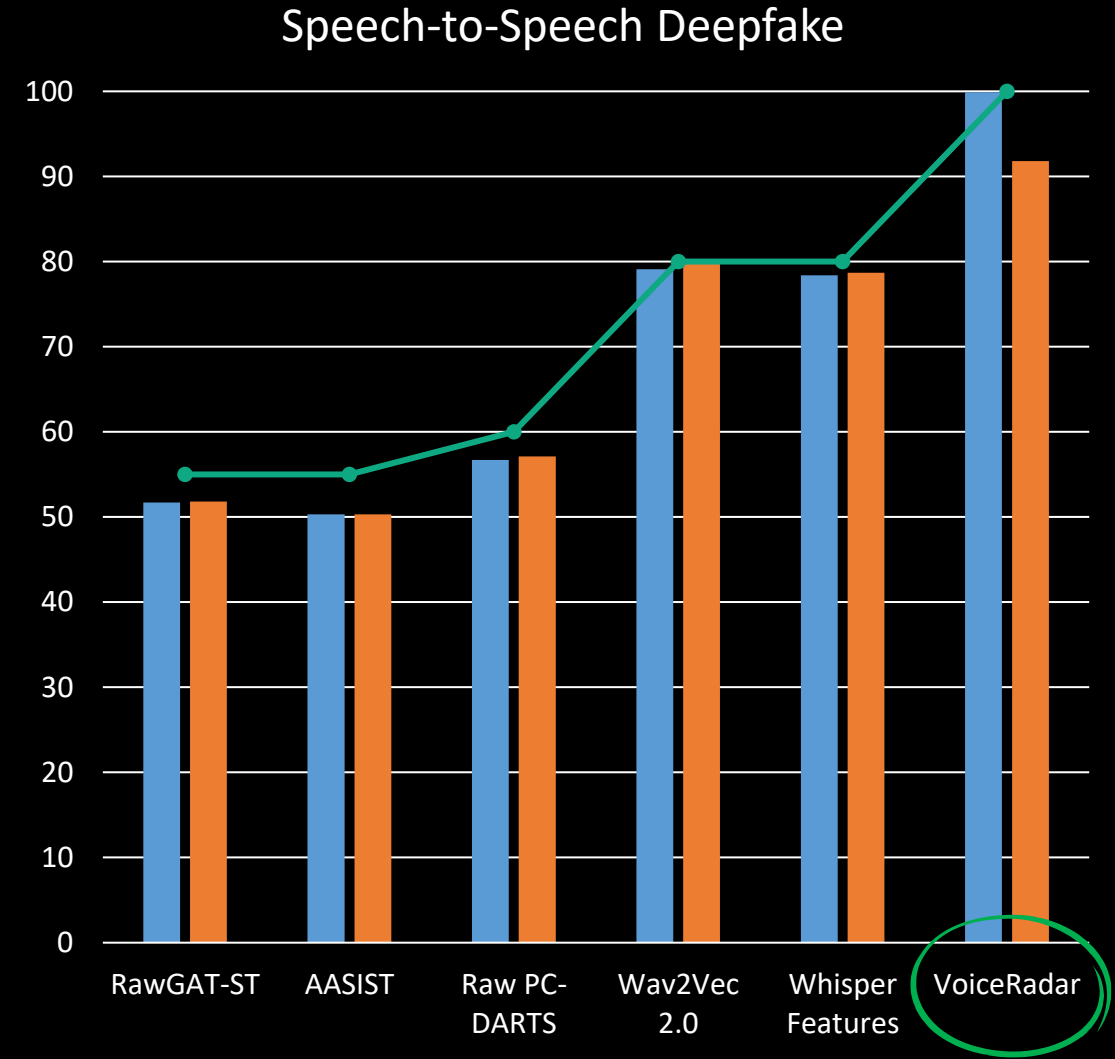
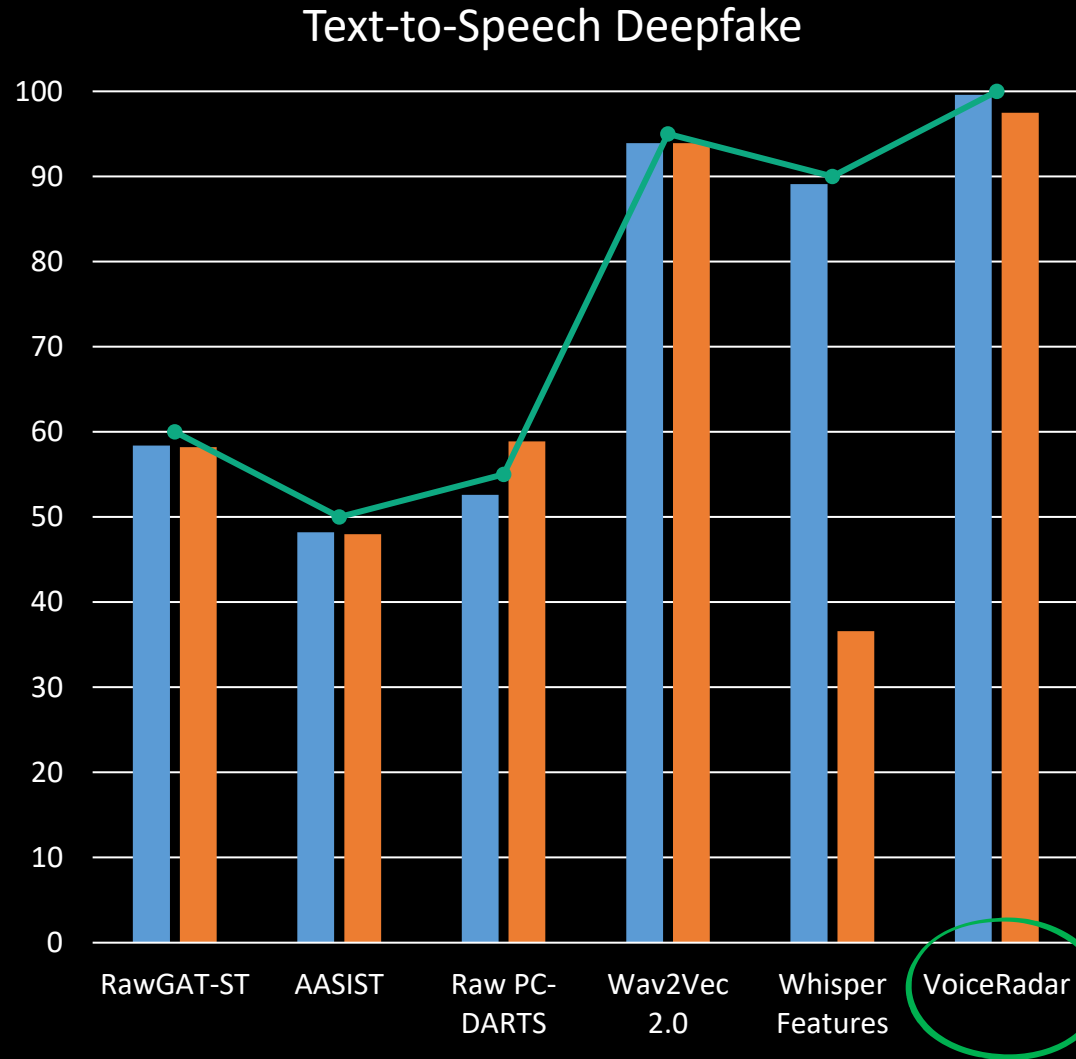
- Cover different distributions of:
Gender, age, and regions



Generated for each speaker different combinations of STS

Dataset	Records
DiffHierVC	145 465
DiffVC	145 465
HierSpeech++	145 483
SpeechT5	145 483
Total	581 896

Comparison of Audio Deepfake Detectors



$$\text{TPR} = \frac{TP}{TP + FN}$$
$$\text{TNR} = \frac{TN}{TN + FP}$$

Conclusion

- Deepfakes are real threats to modern societies
- We addressed existing detectors' limitations
- We proposed VoiceRadar
 - Agnostic of Text-to-Speech or Voice-Conversion

