

Deepfakes – příležitost nebo hrozba?

Kamil Malinka a Anton Firc

malinka@fit.vut.cz, ifirc@fit.vut.cz

Security@FIT





Vzdělávání - Výzkum - Společnost

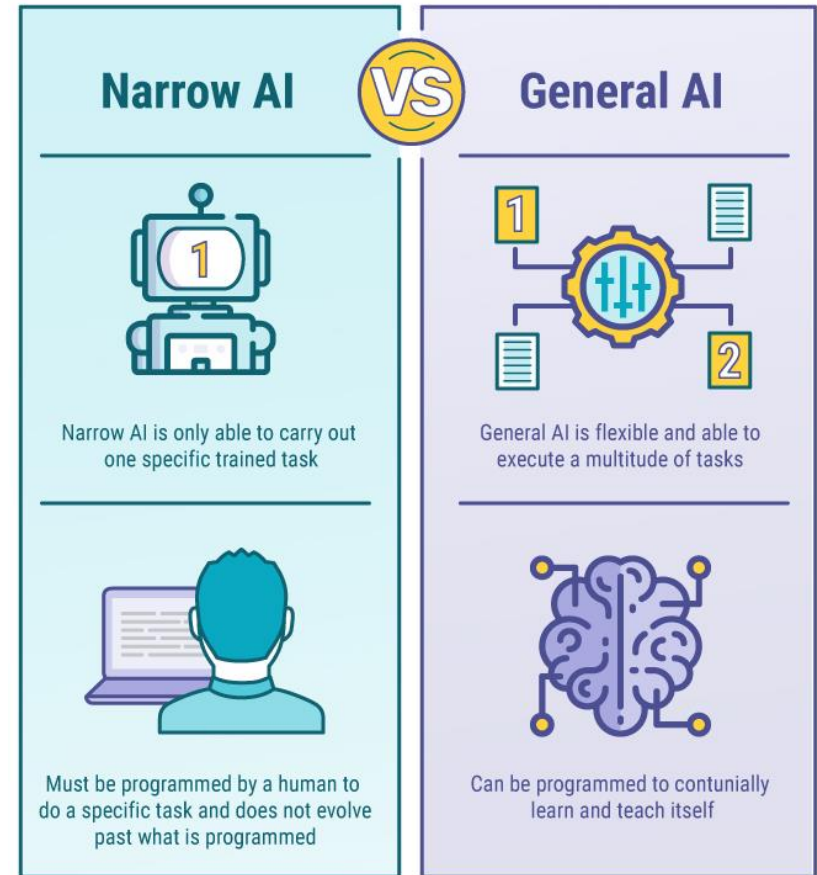
- **Výzkumné směry**

- Bezpečnostní dopady AI
- Blockchainové technologie
- Analýza malware
- IoT technologie
- Použitelná bezpečnost
- Finanční protokoly

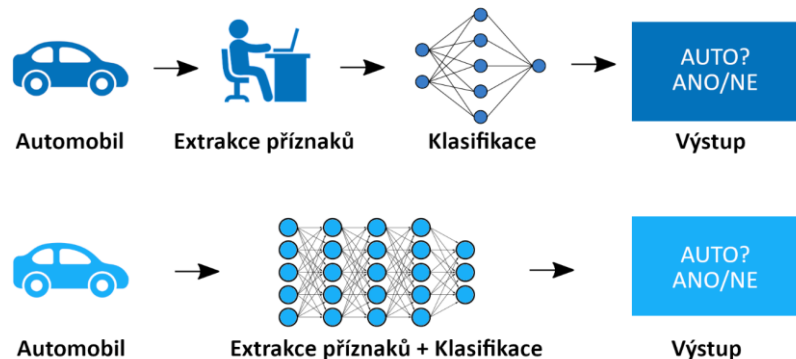
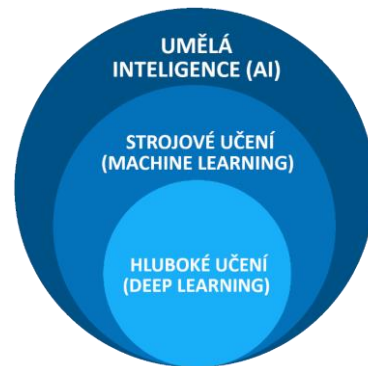


Umělá inteligence: co to vlastně je?

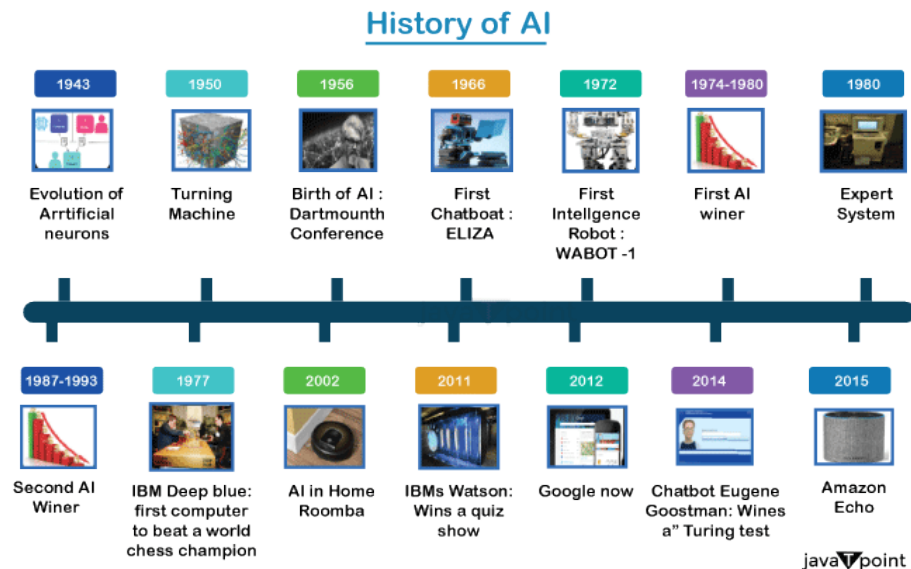
- **AI** (umělá inteligence): Stroj se schopností vyhodnocovat data a činit rozhodnutí stejně dobře nebo lépe než člověk. Rozumí datům a výstupům
- **Narrow AI** (úzce zaměřená AI): AI zaměřená na jeden konkrétní úkon/téma
- **General AI** (všeobecná umělá inteligence): Umělá inteligence schopná jako člověk přizpůsobit své fungování na jakékoliv téma/úkoly



- **Machine Learning** se zabývá tvorbou a studiem statistických algoritmů schopných se učit z dat a aplikovat naučené na nová data, a tím vykonávat úkoly bez nutnosti konkrétních pokynů.
- **Neural network** (neuronová síť) je matematický model pro zpracování dat inspirovaný funkcí lidského mozku.
- **Deep learning** (hluboké učení) je varianta strojového učení využívající neurální síť s mnoha vrstvami k spracování složitých dat a k učení se složitým vzorcům bez nutnosti ručního nastavování

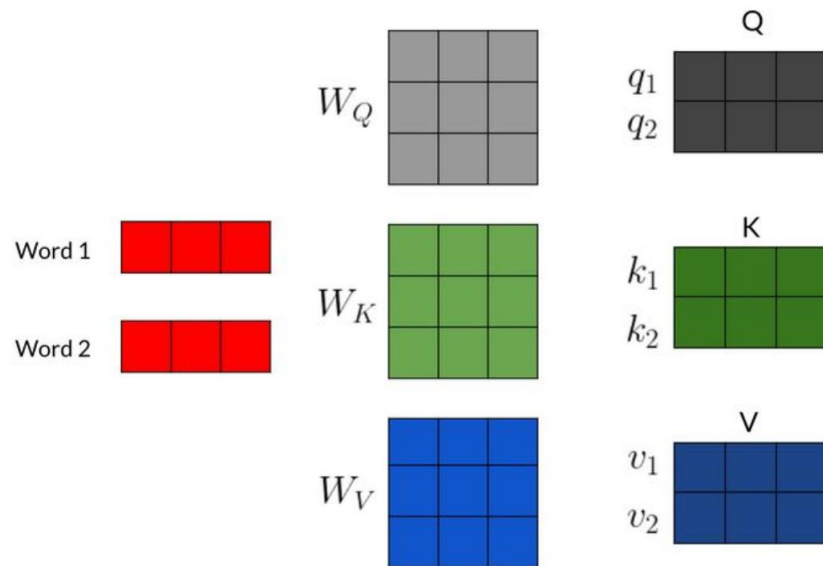


- Teorie neuronových sítí vznikla již v roce 1943
- První praktické pokusy v šedesátých letech a pokroky v osmdesátkách
- Zlom v 21. století
 - vazba na pokroky v hardwaru
- Kdy bude další průlom?



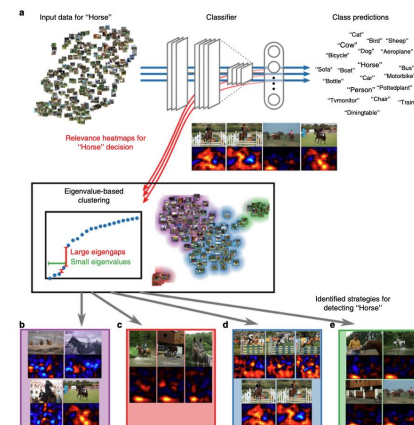
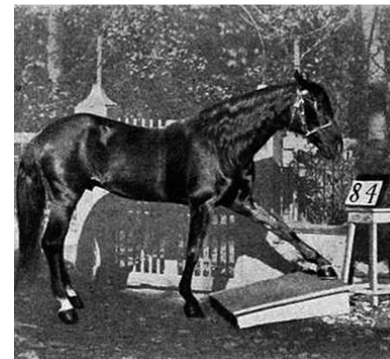
- Pokročilý ML
 - Jak moc ChatGPT reálně rozumí tomu co píše, a kolik jsou jen neskutečná kvanta naučených vzorců?
 - “jednoduché násobení matic”
- To co máme dnes je Narrow AI
- I v případě že máme „jen“ ML – má velký disruptivní potenciál

Here is step by step guide, first you get the Q, K, V matrices:

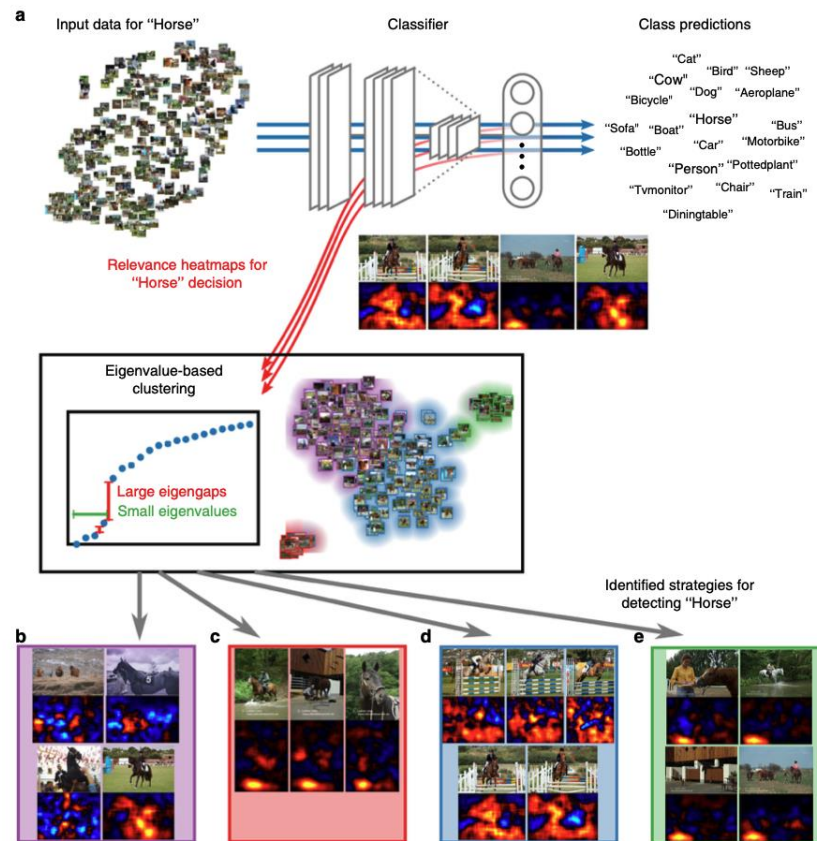


For each word, you multiply it by the corresponding W_Q , W_K , W_V matrices to get the corresponding word embedding. Then you have to calculate scores with those embedding as follows:

- Většina současné AI technologie využívá ML a primárně DL
- Dataset = data ze kterých se AI učí
- Datasetsy jsou stejně, ne-li více důležité než AI algoritmy samotné
- “Clever Hans“ efekt
- Obrovské množství možných biasů, chyb, atd.
- Klíčové know-how



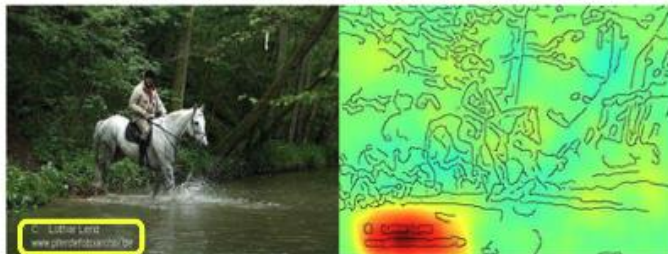
- Většina současné AI technologie využívá ML a primárně DL
- Dataset = data ze kterých se AI učí
- Datasetsy jsou stejně, ne-li více důležité než AI algoritmy samotné
- “Clever Hans“ efekt
- Obrovské množství možných biasů, chyb, atd.
- Klíčové know-how



- Většina současné AI technologie využívá ML na

a

Horse-picture from Pascal VOC data set



Source tag
present



Classified
as horse

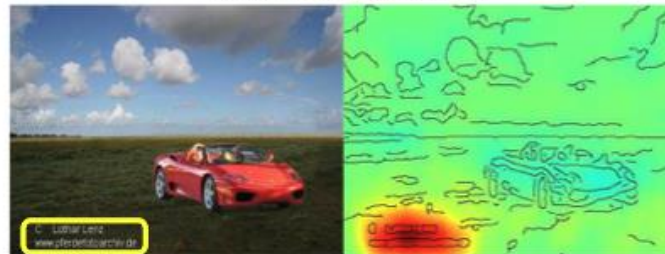


No source
tag present



Not classified
as horse

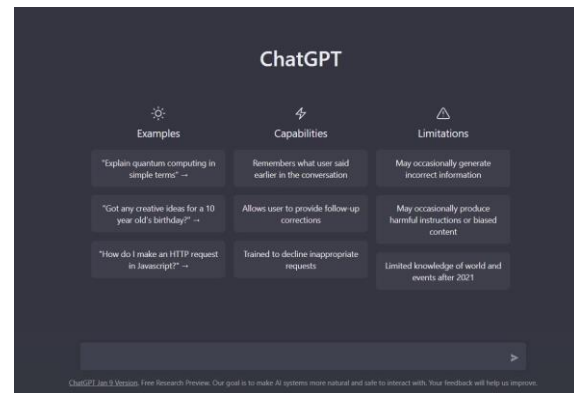
Artificial picture of a car



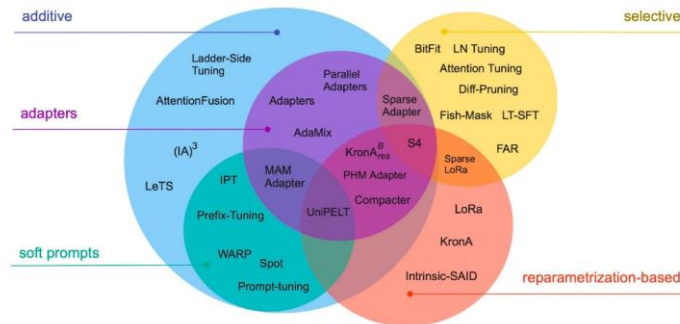
- KNICOVE KNOW-HOW



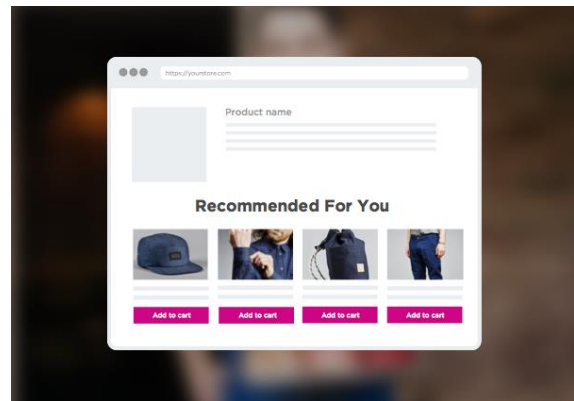
- AI schopná generovat „hmatatelné“ výstupy nebo produkty
- Text, zvuk, obraz, video...
- V současnosti vysoce populární
- Text: ChatGPT a konkurenční chatboti
- Obraz: Dall-E, Midjourney...
- Trend: Multimodální vstupy a výstupy



- Klíč fungování Chatbotů
- Textové datasety převedeny do matematického modelu
- Přiřazuje slovům hodnoty a pravděpodobnosti jejich výskytu poblíž jiných slov
- Na první pohled velmi přesvědčivý, ale pořád velmi limitovaný



- Především AI pro rozhodovací procesy
- Výstupem není produkt, ale rozhodnutí
- Nejsou tolik vidět jako generativní AI, ale neméně důležité
- Samořídící automobily, asistence s diagnózami, algoritmy pro doporučování obsahu, cílené reklamy...



- Současný milník na horizontu je (kvalitní) video
- Velké akční modely? Text-to-game?
- Konečně plně autonomní auta?
- Vazba na hardware
- Boom nebo stagnace?



Umělá inteligence a kyberbezpečnost

- AI promění (a již proměňuje) podobu kybernetické bezpečnosti
 - Ve všech rovinách
- AI slouží a bude sloužit útočníkům
 - zároveň ale i obráncům
- Závody v AI zbrojení?
- Morální a etické limity „těch hodných“
- AI se v neposlední řadě sama stává cílem jak AI tak ne-AI útoků

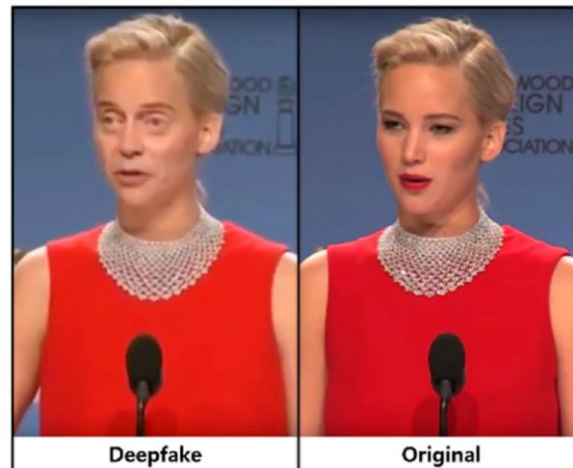
- Automatické hledání zranitelností
- Monitorování síťového provozu
- C2 (chytrý autonomní malware)
- Automatické shánění podkladů pro social engineering
- Řízení masivních a chytrých botnetů

- Psaní dezinformací
- Generování dezinformačních obrázků
- Oboje využitelné pro phishing
- Deepfakes (videa a hlas)
- Vishing
- Psaní malwaru
 - Zatím málo sofistikovaný, ale dostupný komukoliv
 - Dedikované „bad guy“ modely

ChatGPT Is Pretty Good at Writing Malware, It Turns Out

The trendy new chatbot has many skills, and one of them is writing "polymorphic" malware that will destroy your computer.

By Lucas Ropek Published January 20, 2023 | Comments (4) | Alerts



- Bojovat s AI pomocí AI
- Vše co AI dělá dokáže AI mitigovat (...teoreticky)
- Chytré IDS
- Chytrý antimalware
- Automatické psaní oprav?
- AI na detekci deepfakes
- Chytré spam filtry
- Útočníci vždy o krok napřed
- Morální a právní limity obránců

- Podíl AI služeb a produktů raketově roste
- Točí se v nich obrovské množství peněz
- Masivně na nich rovněž roste závislost společnosti
- Nevyhnutelné útoky proti AI (ať už s AI nebo bez)
- V menší míře už se dějí, a bude jich přibývat



Forbes

FORBES > INNOVATION > CYBERSECURITY

ChatGPT Down As Anonymous Sudan Hackers Claim Responsibility

Davey Winder Senior Contributor @
*Veteran cybersecurity and tech analyst, journalist,
hacker, author*

Follow

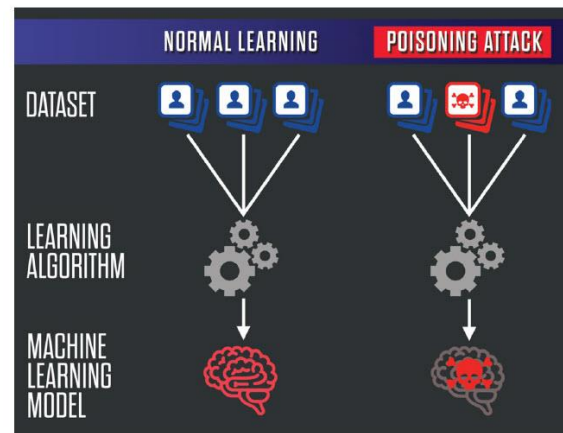


Nov 10, 2023, 08:39am EST

- DDoS – Tradiční denial of service bude v nejbližší době hlavním nástrojem
- Ransomware – může zasáhnout i AI, velmi lákavé cíle
- Jailbreak – přinutit chatbota k podání zakázaných informací



- Narušení „black boxu“ – potenciálně těžce odhalitelné
- Nemusí ale mířit na AI přímo
- Otrava datasetů
- Nazpět k příkladu s koňmi
- Co kdyby AI letištní skener skrze otrávené datasety nepoznal zbraň v batohu?
- Proto potřeba chránit jak AI tak datasety



Deepfakes

Název **Nástroje boje proti hlasovým DeepFakes (VB02000060)**

Tools To Combat Voice DeepFakes



Realizace 01/2024 - 12/2026



Realizuje FIT VUT v Brně (Speech@FIT + Security@FIT) + Phonexia s.r.o.

Uživatel Policejní prezidium ČR

Poskytovatel MV ČR - Program bezpečnostního výzkumu ČR 2021 – 2026 (SECTECH)

Veřejný link <https://www.fit.vut.cz/research/project/1762>

#nePINdej!

***|
KYBERTEST



*** |

KYBERTEST



...|
KYBERTEST

#nePINdej!
Nikdy nikomu
nedávej
svoje hesla!

Jasný!
E-šmejdi mě
nedostanou!
#nePINdám!

**E-šmejdi vás
chtějí okrást!**

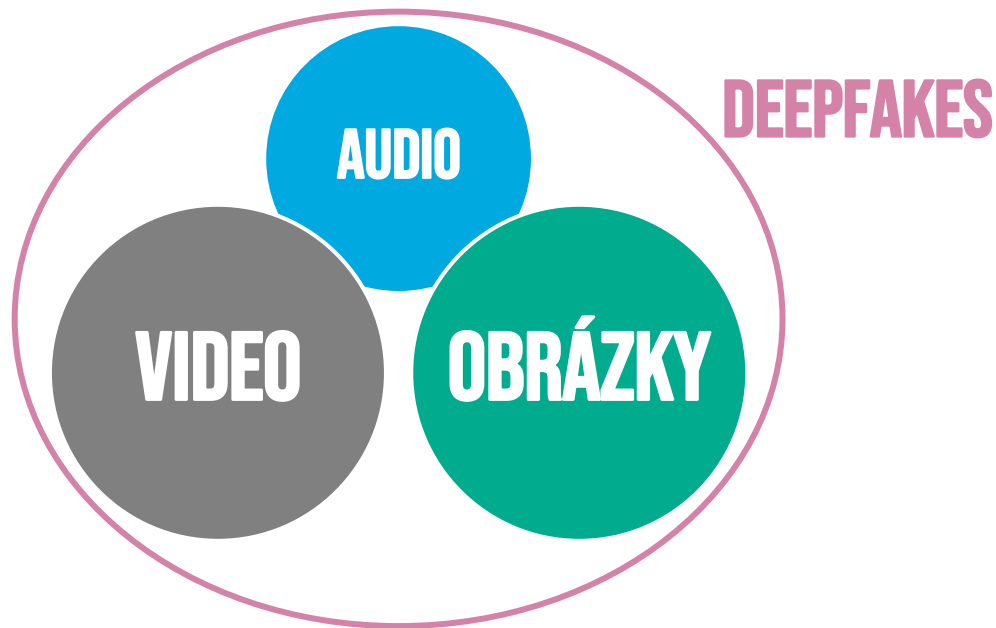
**Naučte se,
jak jim nenaletět!**

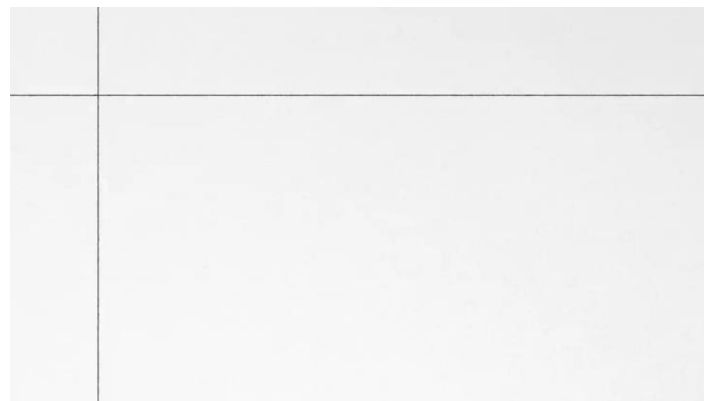
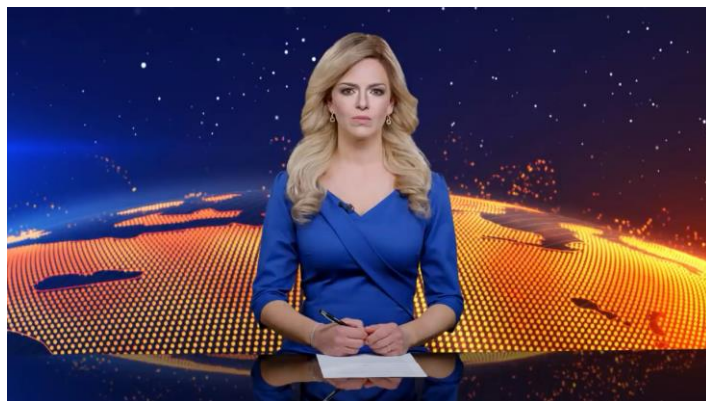


www.kybertest.cz

- Motivace a reálné útoky
- Trocha teorie k deepfakes
- Metody ochrany
- Detaily k výzkumným oblastem (když vyjde čas...)
 - Útok na hlasovou biometrii
 - Útoky na obličejovou biometrii
 - Lidské rozpoznávání deepfakes

- “deep learning” + “fake”
- Podmnožina syntetických médií
- Zobrazují události které se nikdy nestaly
- Výrazné zvýšení kvality a dostupnosti v posledních letech







**NOT HIS
REAL
VOICE**

**TOP GUN
MAVERICK**







Demonstrační deepfake ukázka vytvořená bez souhlasu vyobrazené osoby.

Simply upload a photo or a video (limited to 50 MB), the demo will tell you whether it is a 'fake'!

pavel.mp4

PERFORM DEEPPAKE DETECTION

BWS Deepfake Detection for 'pavel.mp4' says:

This was a fake video!



Demonstrační deepfake ukázka vytvořená bez souhlasu vyobrazené osoby.

Simply upload a photo or a video (limited to 50 MB), the demo will tell you whether it is a 'fake'!

prezident.mp4

PERFORM DEEPAKE DETECTION

BWS Deepfake Detection for 'prezident.mp4' says:

The video has been recorded from a live person!

HN.cz > Archiv
[N]

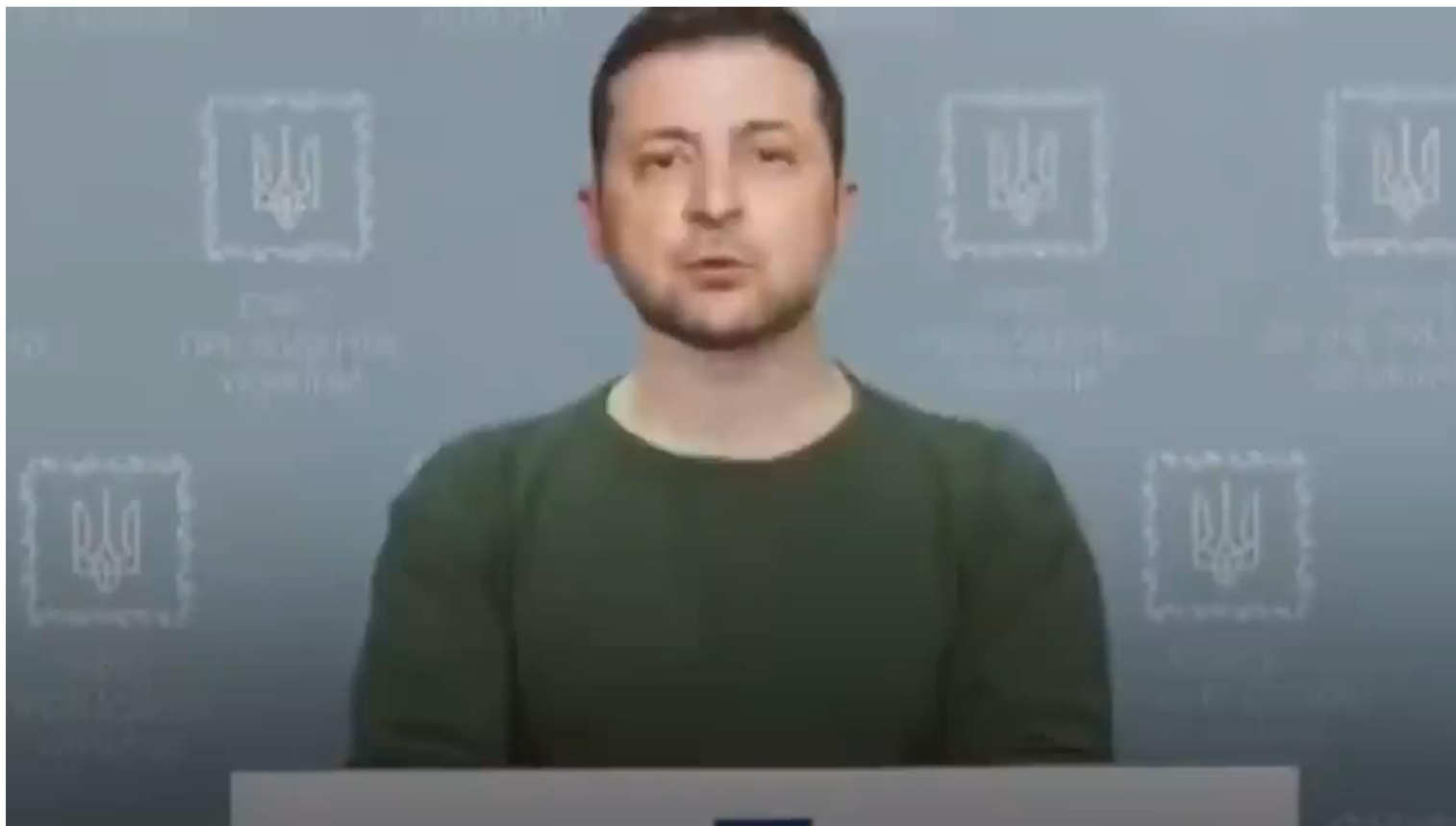
Man created 'deepfake' voice of his classmate using their voice. Miliardový e-shop se stal terčem deepfake útoku. Zaměstnanec si čtvrt hodiny volal s videokopíí šéfa GymBeamu

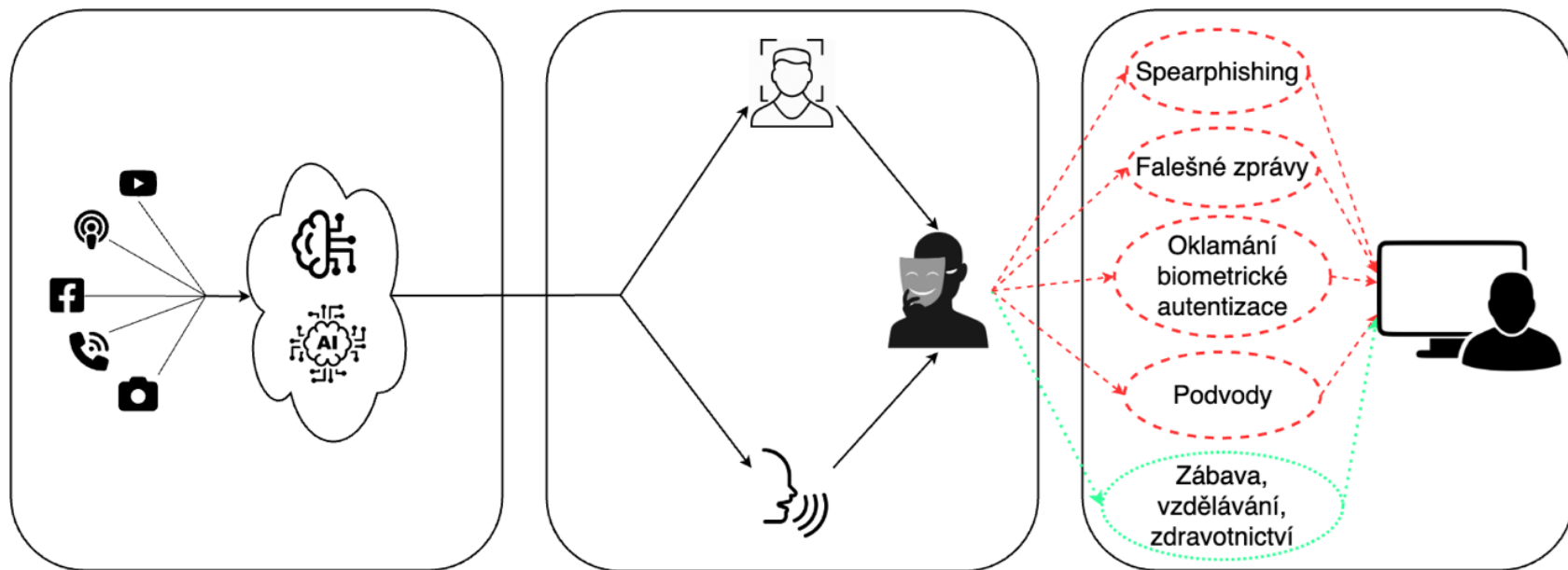
Fake Biden robocall tells voters to skip New Hampshire primary electionidents used in Russia-

23 January 2024

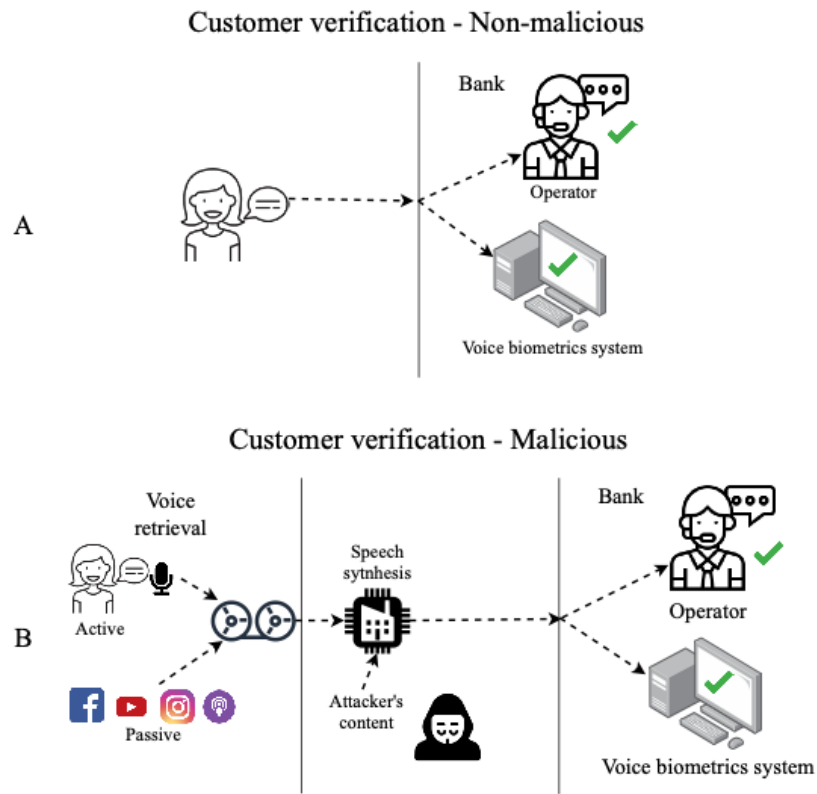
Startup Shocked When 4Chan Person who has spoken in Vienna AI Immediately Abuses Its Voice-Cloning

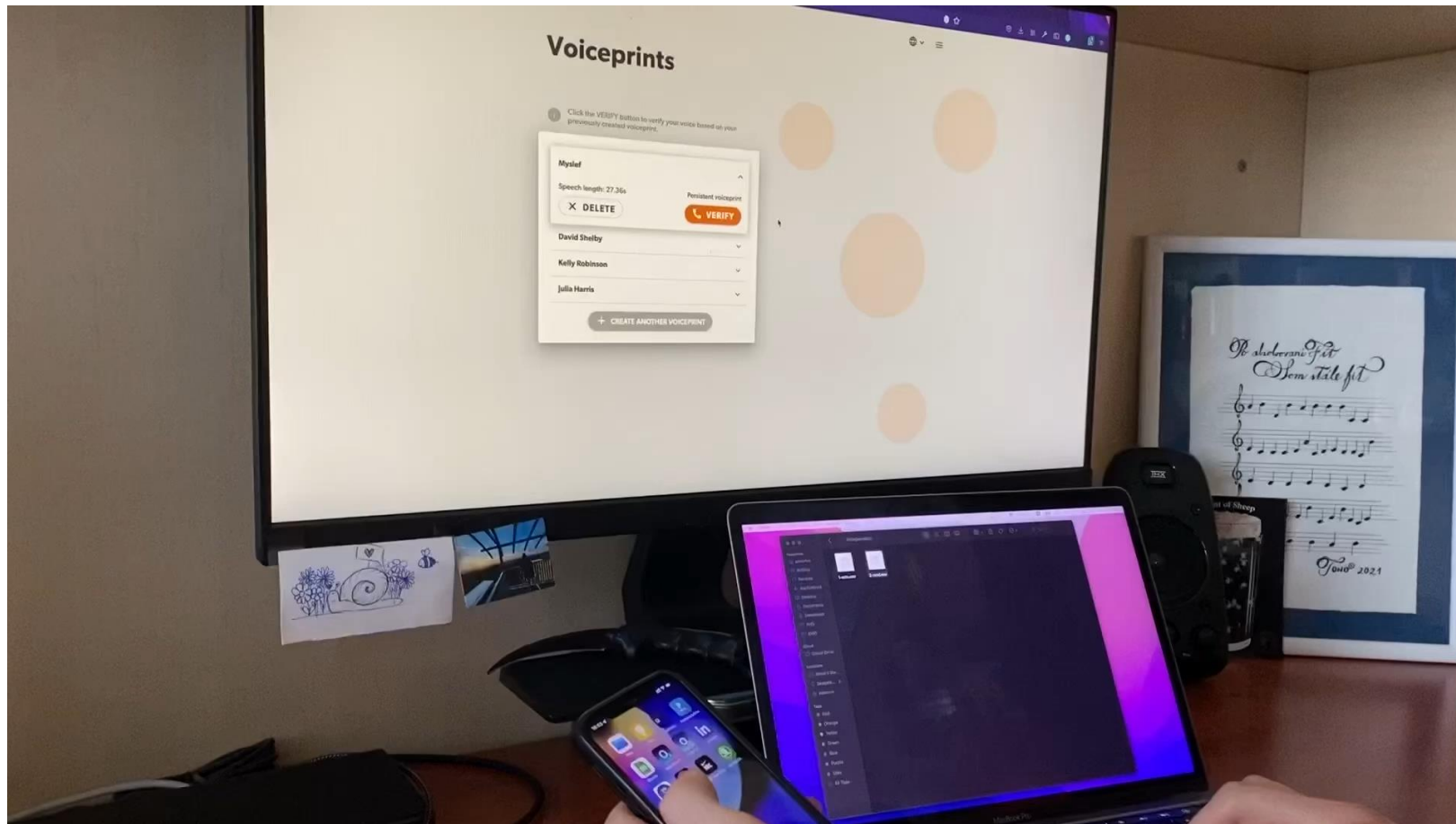
Deepfaked Voice Enabled \$35 Million Bank Heist in 2020





- Dva faktory:
 - Hlasová biometrie
 - Lidský operátor
- 3 hlasové syntetizátory + 2 hlasové biometrie







Face-swap v reálné čase ↗

- Lidé obecně nedokážou rozlišit mezi skutečnou a falešnou řečí.
 - Operátor není varován před vystavením deepfakes
- Pečlivě připravené deepfake nahrávky lidi oklamou
- Větší důraz je kladen na obsah konverzace, aby bylo možné identifikovat autenticitu řeči

96.8% lidí **neodhalí** deepfake v běžné konverzaci


NO DEEPPFAKE DETECTED
New Scan



Name:

cXFL-P0045o

User

Source

2021-10-11 09:34:49 UTC

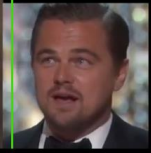
-42 minutes ago

Size:

1.3 MB

DETAILS

Deepware aims to give an opinion about the scanned video and is not responsible for the result. As Deepware Scanner is still in beta, the results should not be treated as an absolute truth or evidence.

Model Results

Avatarify: NO DEEPPFAKE DETECTED(3%)

Deepware: NO DEEPPFAKE DETECTED(2%)

Seferbekov: NO DEEPPFAKE DETECTED(4%)

Ensemble: NO DEEPPFAKE DETECTED(2%)

Video

Duration: 15 sec

Resolution: 1280 x 720

Frame Rate: 29.97 fps

Codec: h264

Audio

Duration: 15 sec

Channel: stereo

Sample Rate: 44 kHz

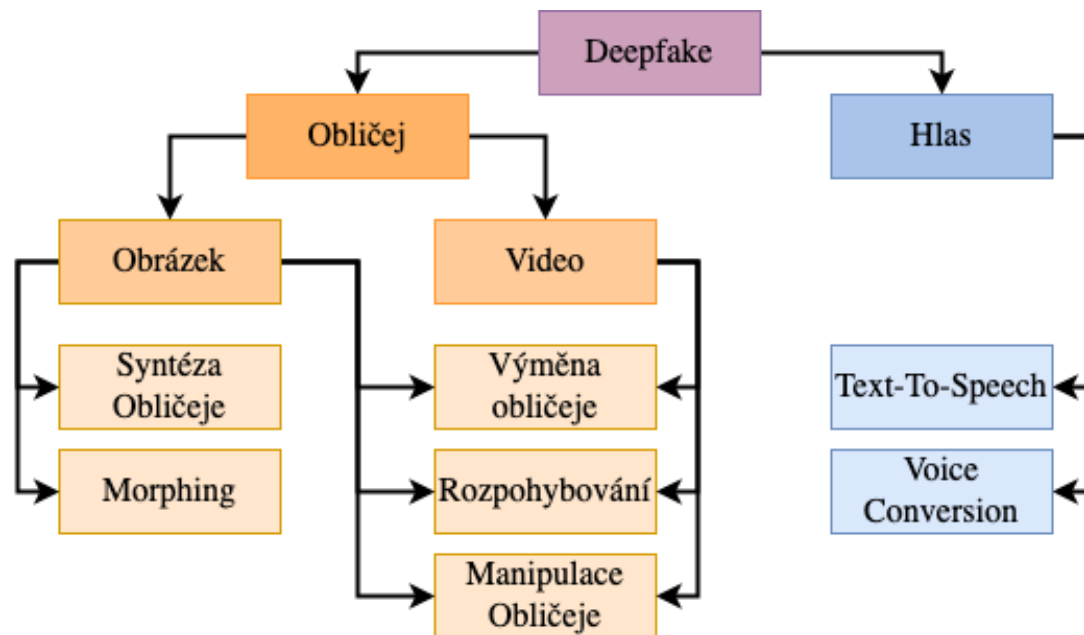
Codec: aac



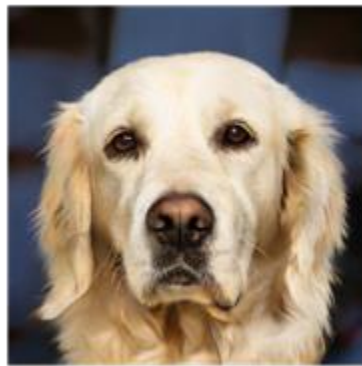
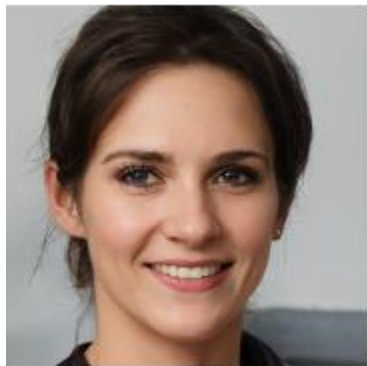
- Blížíme se do stavu tvorby deepfakes v reálném čase (spíše už tam jsme..)
 - Např. filtry na MS Teams
- Deepfakes přinášejí nové rizika pro společnost
 - Biometrická autentizace
 - Social engineering
- Detekce deepfakes potřebná pro ochranu systémů
- Lidské rozpoznávání deepfakes
- Zvyšující se kvalita a dostupnost
- S rozvojem jazykových modelů (ChatGPT apod.) se blíží riziko automatizace útoků

- Útočník 1 (vše automatizovaně a se širokým pokrytím):
 - Stáhne fotky ze sociálních sítí
 - Spojí je s osobními údaji, kontakty atp.
 - Vytvoří pomocí face swap nahé fotky
 - Začne vydírat a hrozit zveřejněním (rozešle výhružku i s ukázkou)
- Útočník 2 (vše automatizovaně a se širokým pokrytím):
 - Vishingové útoky (chatbot vydávající se za osobu blízkou v nesnázích)

Jaké existují deepfakes?



- Syntéza ne-existujících objektů
 - Lidi, zvířata, scenérie, auta...
- Trénink nad velkým množstvím obrázků stejné kategorie
- Náhodné generování z naučeného latentního prostoru





1



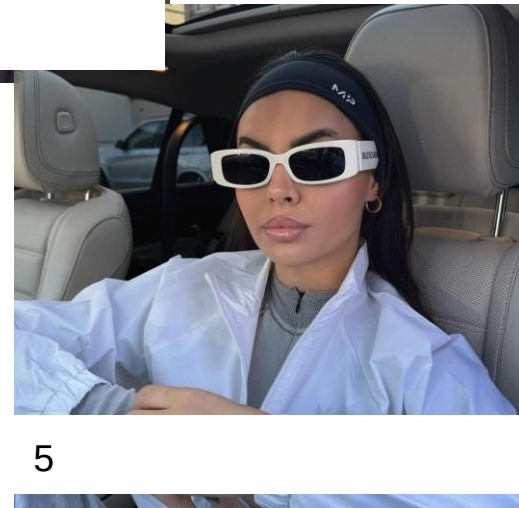
2



3



4



5



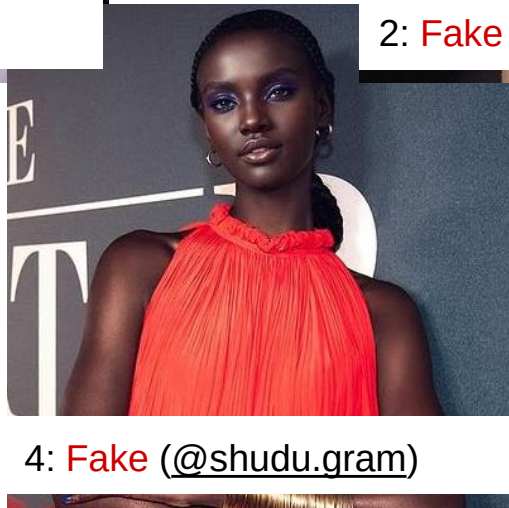
1: **Real** (@beth.sal)



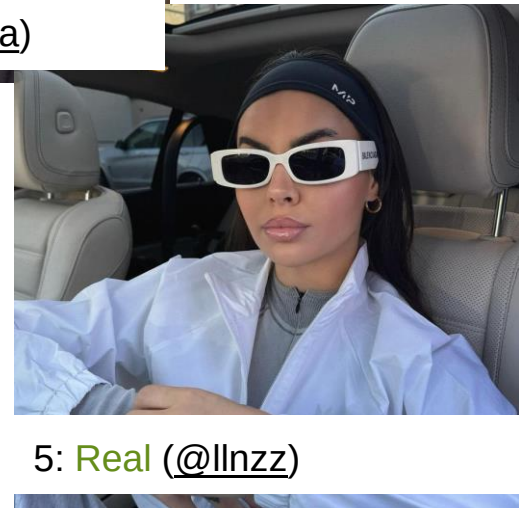
2: **Fake** (@fit_aitana)



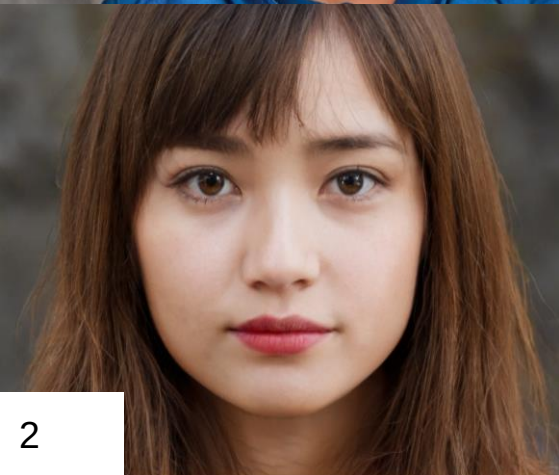
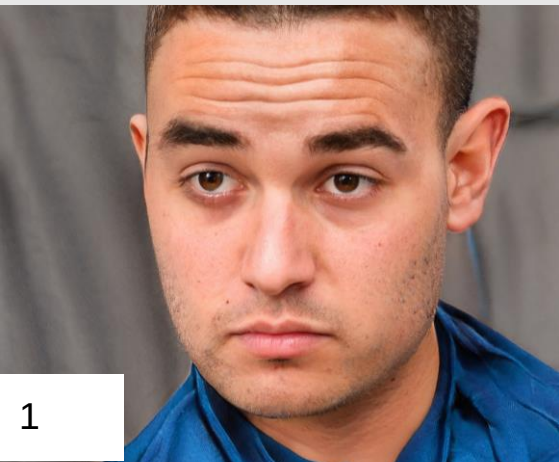
3: **Fake** (@bejby.blue)



4: **Fake** (@shudu.gram)

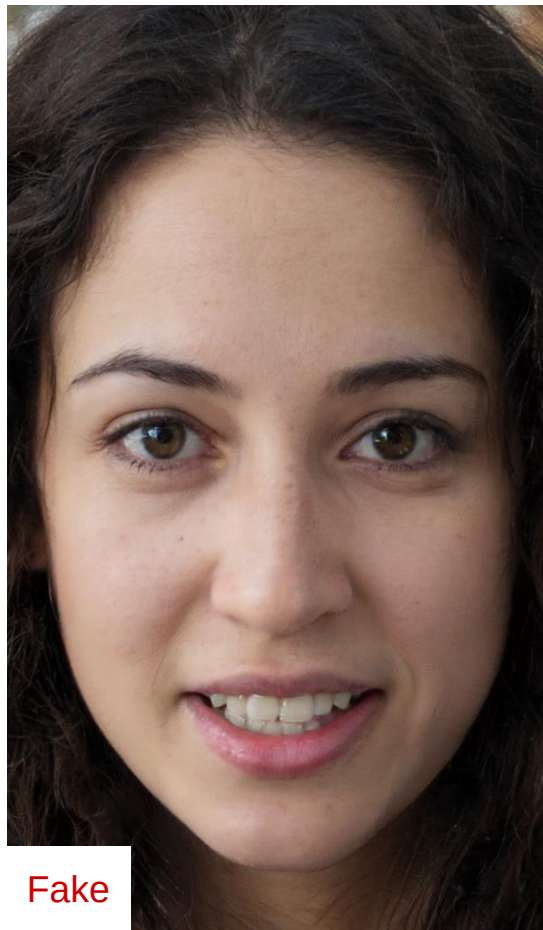


5: **Real** (@llnzz)





Fake



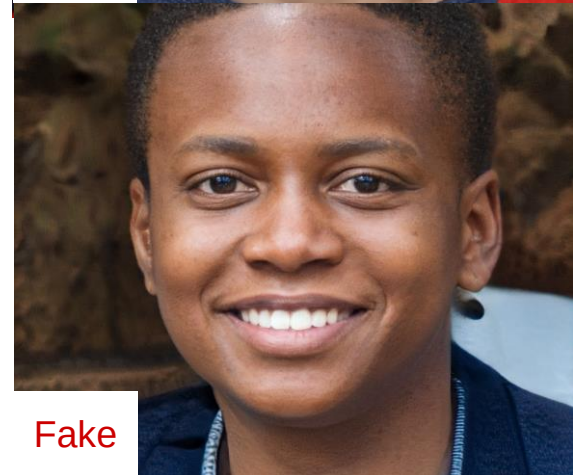
Fake



Fake



Fake



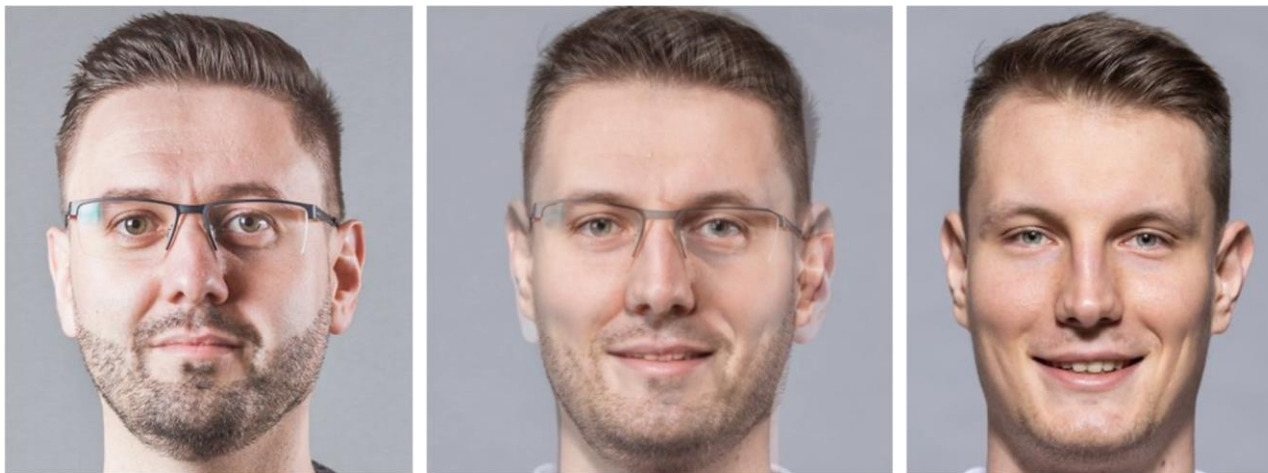
Fake

- Boti na sociálních sítích
 - Syntetické profilové obrázky = náročná detekce
 - Falešné Twitter účty šířící dezinformace v Belgii [1]
- Negativní tvrzení o identitě
 - “nemám přístup do systému, ale měl bych“



[1] Graphika Team, Fake Cluster Boosts Huawei, 2022.
<https://graphika.com/reports/fake-cluster-boosts-huawei>.

- Morphing = efekt změny z jednoho obrázku na druhý
- 2 obličejů -> jeden podobný oběma
- Nebezpečí pro automatizované hraniční kontroly identity



- Problém pro obličejové biometrie
 - Automatické hraniční kontroly
- Dva lidi, jeden občanský průkaz (pas)
- Incidenty hlásí Slovinská policie [1]
 - 40 incidentů
 - Zneužívání Slovinských pasů ne-EU občany pro cesty do Kanady



<https://www.schengenvisainfo.com/news/estonia-implements-new-automated-border-control-gates-for-travellers/>

[1] Chris Burt, Morphing attack detection for face biometric spoofs needs more generalization, datasets: Biometric Update.
<https://www.biometricupdate.com/202207/morphing-attack-detection-for-face-biometric-spoofs-needs-more-generalization-datasets>, 2022.

- Jeden dokument = dva "právoplatní" majitelé



=



Similarity rate: 88%



=



Similarity rate: 78%



≠



Similarity rate: 0%

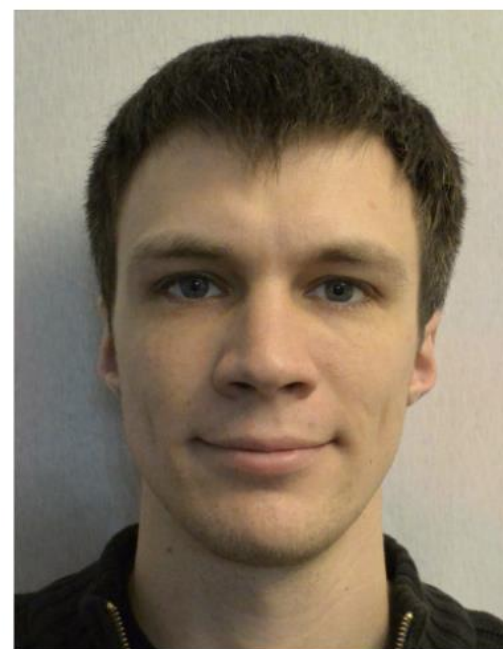
[1] <https://faceapi.regulaforensics.com/>



(a) Subject 1

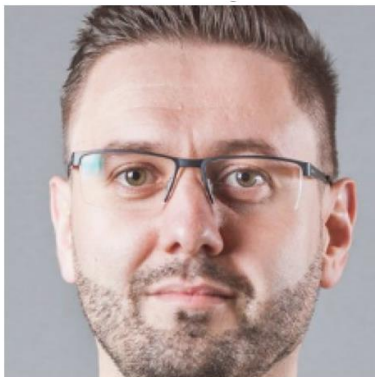


(b) Morph

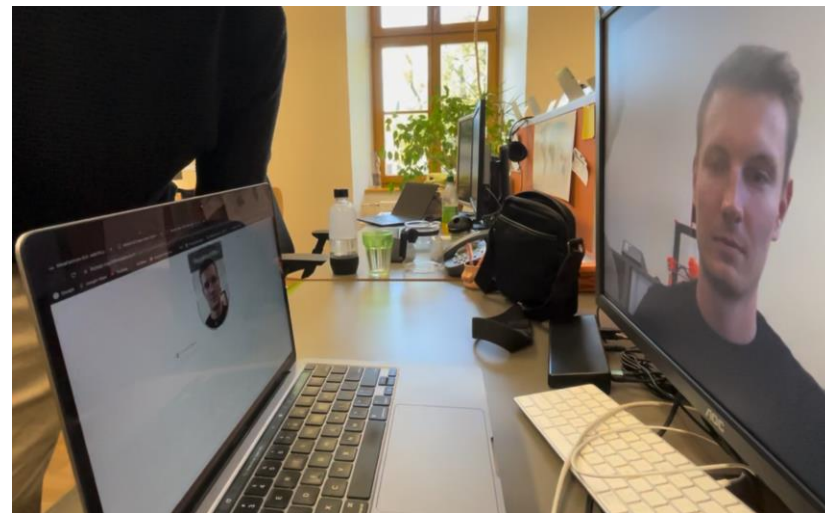


(c) Subject 2

- Obličej ze zdrojové fotografie je přenesen na hlavu v cílové fotografii
- Obrázky a video



- Útok na biometrické systémy
- Krádež identity
- Pomluva
- Manipulace s důkazy
- Úprava pornografického materiálu
 - Sdílení upravených fotek s obličejí studentek v USA [1]
 - 98 % deepfake videí má pornografický obsah [2]
 - 99 % z nich je zaměřeno na ženy.



[1] Thaler, S. (2023, November 2). AI-generated nude images of girls at NJ high school trigger police probe: "I am terrified." New York Post. <https://nypost.com/2023/11/02/news/ai-generated-nudes-of-girls-at-nj-high-school-trigger-police-probe/>

[2] Security Hero. (2023). STATE OF DEEPPAKES 2023. <https://www.securityhero.io/state-of-deepfakes/>

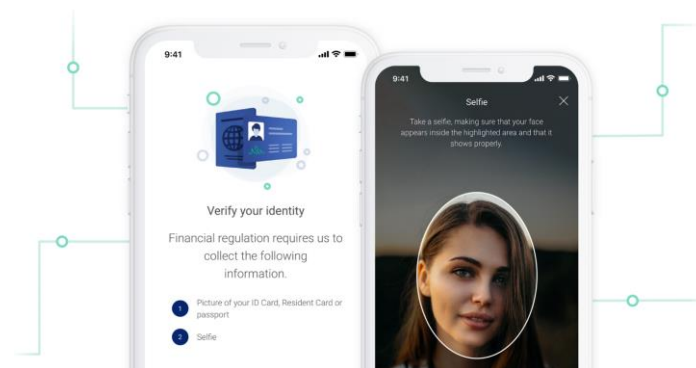


Face-swap v reálném čase ↑

- Foto-realistická re-animace obrázku podle zdrojového videa
- Generování vizuální informace a synchronizace audia a videa



- Krádež identity
- Falešné zprávy
- Know Your Customer proces
 - Ověření identity v aplikaci na základě fotky dokladu a nasnímání obličeje
 - Možnost oklamat test živosti
- Únik na daních v hodnotě 76 mil. USD [1]
 - Nahrání falešných identit do daňového portálu v Číně



<https://medium.com/2gether/this-is-why-the-kyc-know-your-customer-is-so-important-for-you-df3f5f1f4e29>

[1] Masha Borak, Tax Scammers Hack Government-Run Facial Recognition System, 2021. <https://www.scmp.com/tech/tech-trends/article/3127645/chinese-government-run-facial-recognition-system-hacked-tax>.

- Technika úpravy vybrané části obličeje v obrázku nebo na videu
- Přidání nebo odstranění vlasů, brýlí nebo tetování
- Úprava atributů v latentním prostoru a zpětná generace obličeje



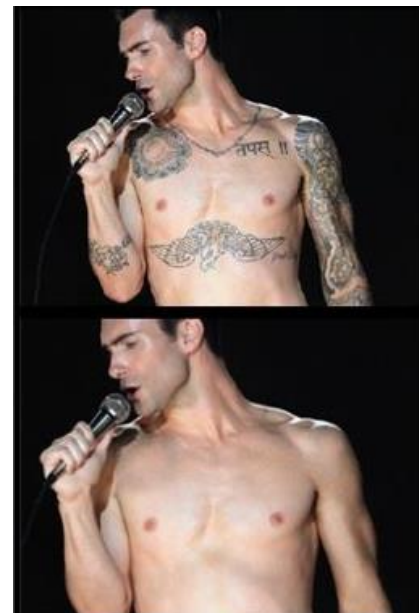
- Úprava atributů obličejů tak, aby osoba nebyla rozpoznána biometrickým systémem
- Manipulace důkazů u soudu



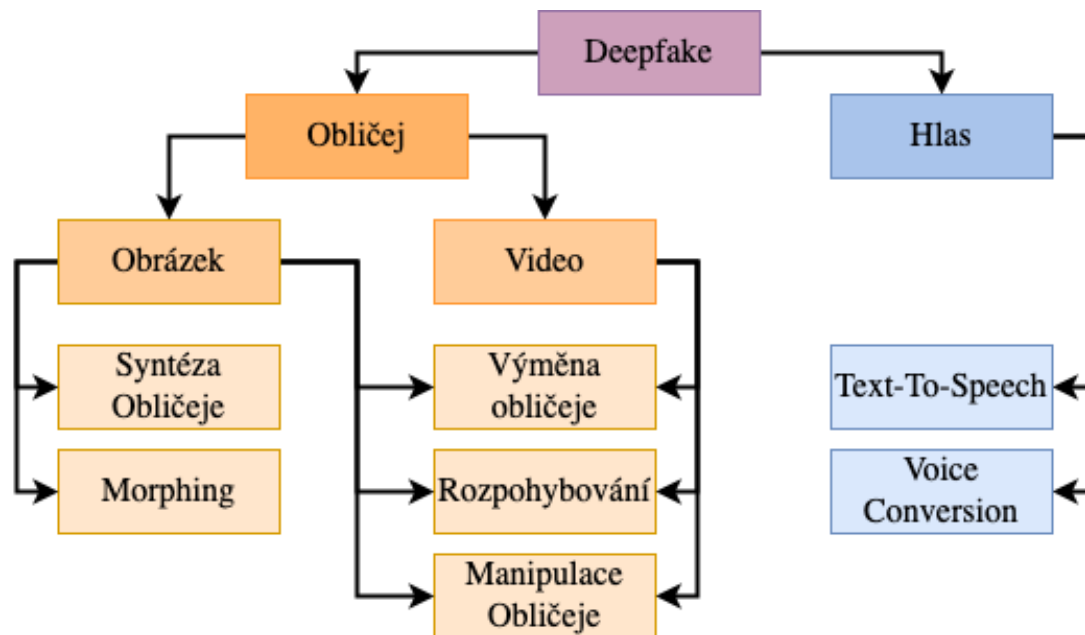
Původní obličej



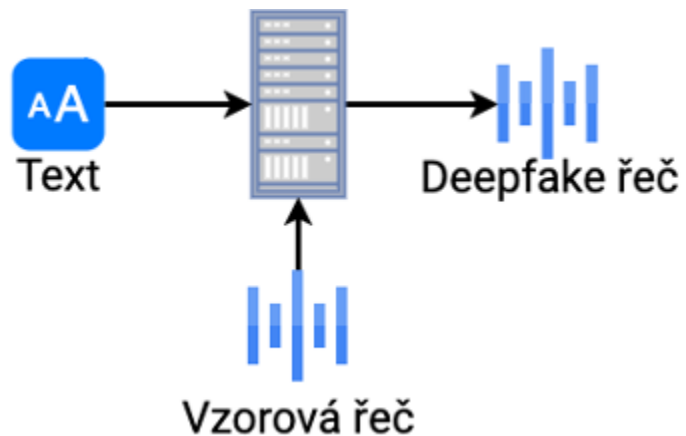
Upravený obličej



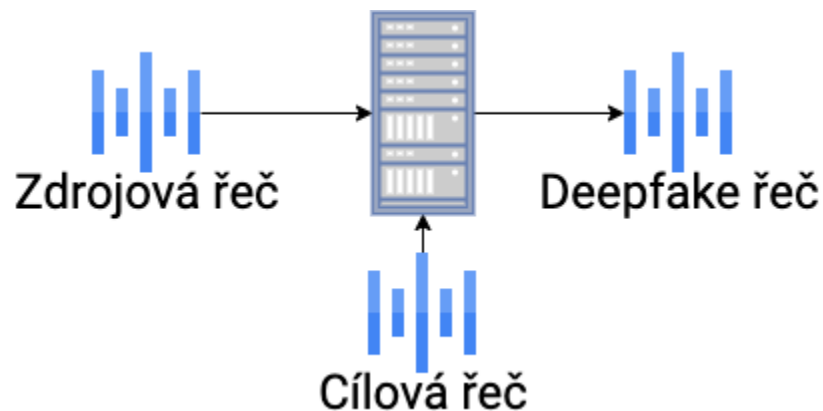
<https://github.com/vijishmadhavan/SkinDeep>



- Převod textu na řeč v hlase ze vzorové nahrávky



- Změna hlasu v zdrojové nahrávce na hlas z cílové nahrávky při zachování obsahu zdrojové nahrávky



Která nahrávka je pravá a která deepfake?

Můžeme?



Fake



Fake



Fake



Fake

- Vhishing – voice phishing
- Krádež identity
- Pomluva
- Podvody
 - Krádež 250 tis. USD [1]
 - Krádež 36 mil. USD [2]

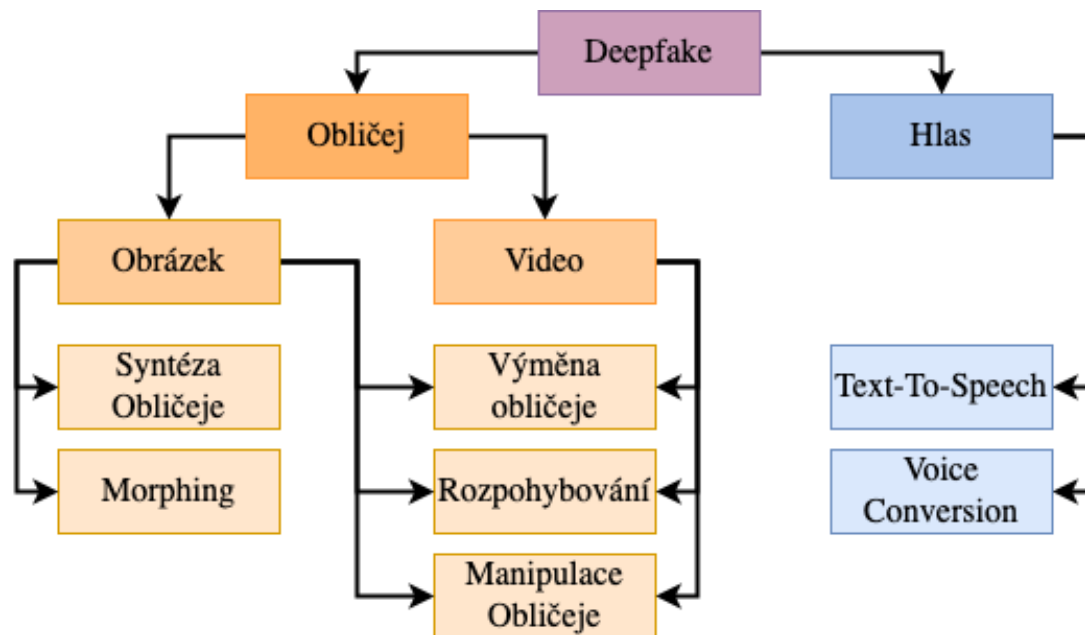


<https://whatismyipaddress.com/vishing>

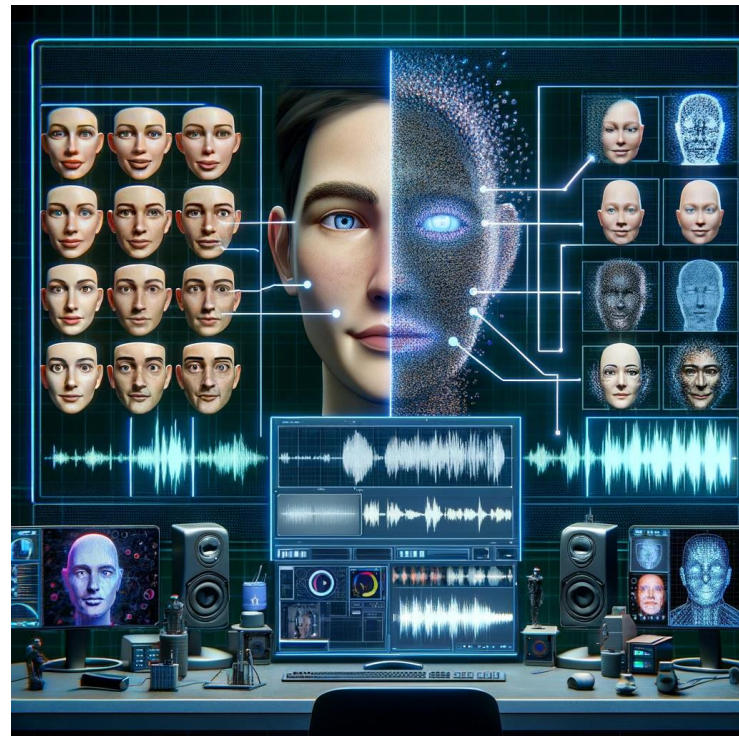


[1] Lindsey O'Donnell, CEO 'Deep Fake' Swindles Company Out of \$243K. <https://threatpost.com/deep-fake-of-ceos-voice-swindles-company-out-of-243k/147982/>, 2019.

[2] Thomas Brewster, Fraudsters Cloned Company Director's Voice in \$35 Million Bank Heist, Police Find, 2021. <https://www.forbes.com/sites/thomasbrewster/2021/10/14/huge-bank-fraud-uses-deep-fake-voice-tech-to-steal-millions/>.



- Spojení obrazových a hlasových deepfakes
 - Např. Rozpohybování + změna hlasu
- Vysoká přesvědčivost



- Nové možnosti útoků
 - Online video konference
- Vyšší přesvědčivost
- Šíření fake news
- Krádež identity
- Podvody
 - Krádež \$25 mil. za použití vícero deepfake identit [1]
 - GymBeam online meeting [2]



[1] Chen, H. (2024, Feb 4). Finance worker pays out \$25 million after video call with deepfake 'chief financial officer'. CNN. <https://edition.cnn.com/2024/02/04/asia/deepfake-cfo-scam-hong-kong-intl-hnk/index.html>

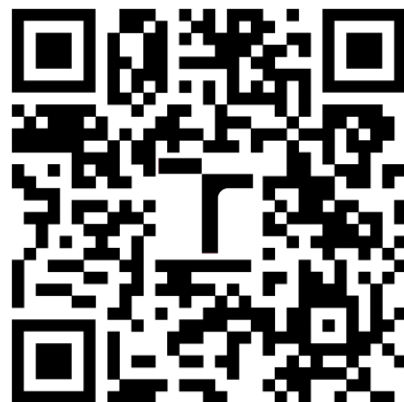
[2] Massayová, V. (2023, August 13). Podvodníci si trúfli na známou slovenskú firmu: V deepfake videohovore sa vydávali za šéfa. Startitup.sk. <https://www.startitup.sk/podvodnici-si-trufli-na-znamu-slovensku-firmu-v-deepfake-videohovore-sa-vydavali-za-sefa/>

Review article

Deepfakes as a threat to a speaker and facial recognition: An overview of tools and attack vectors

Anton Firc^{*}, Kamil Malinka, Petr Hanáček

Brno University of Technology, Božetěchova 2, Brno, 612 00, Czech Republic



ARTICLE INFO

Keywords:

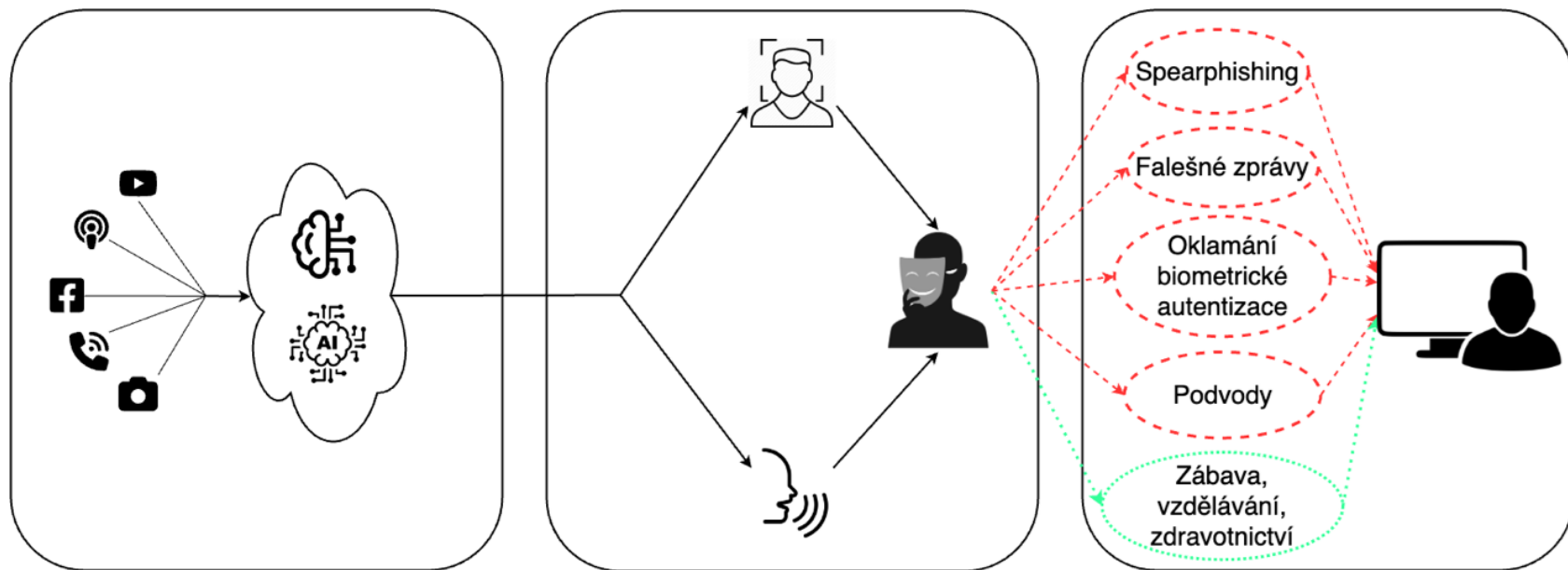
Face deepfakes
Speech deepfakes
Biometrics systems
Facial recognition
Speaker recognition
Deepfake detection
Cybersecurity

ABSTRACT

Deepfakes present an emerging threat in cyberspace. Recent developments in machine learning make deepfakes highly believable, and very difficult to differentiate between what is real and what is fake. Not only humans but also machines struggle to identify deepfakes. Current speaker and facial recognition systems might be easily fooled by carefully prepared synthetic media – deepfakes. We provide a detailed overview of the state-of-the-art deepfake creation and detection methods for selected visual and audio domains. In contrast to other deepfake surveys, we focus on the threats that deepfakes represent to biometrics systems (e.g., spoofing). We discuss both facial and speech deepfakes, and for each domain, we define deepfake categories and their differences. For each deepfake category, we provide an overview of available tools for creation, datasets, and detection methods. Our main contribution is a definition of attack vectors concerning the differences between categories and reported real-world attacks to evaluate each category's threats to selected categories of biometrics systems.

Metody ochrany





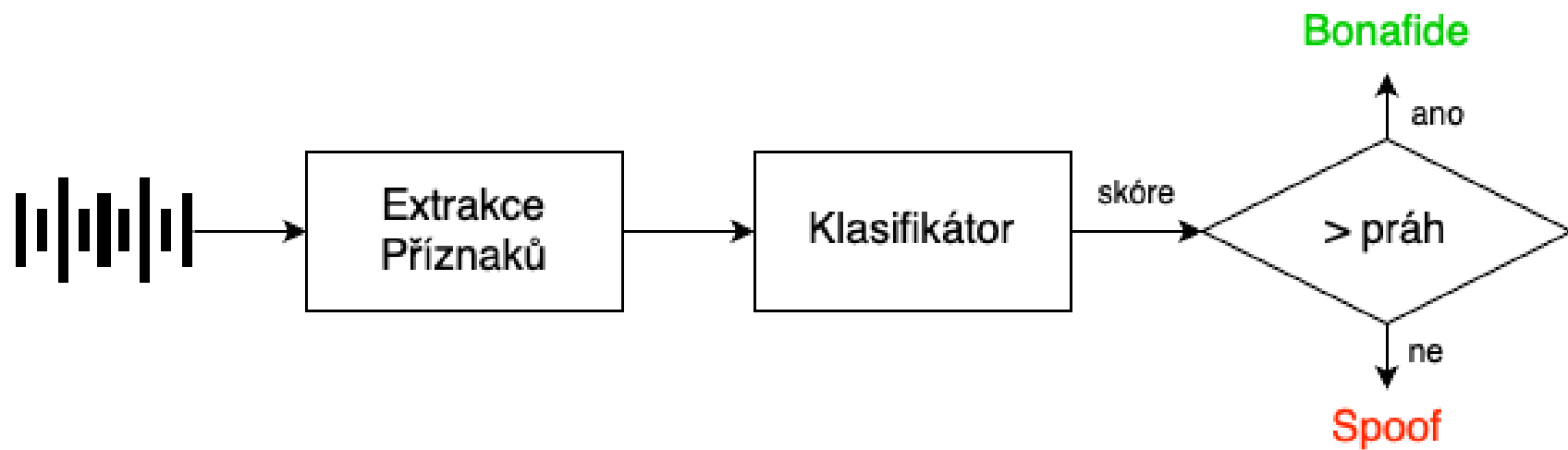
- Tři oblasti pro implementaci ochran
 - Získávání dat
 - Watermarking
 - Právní regulace
 - Tvorba deepfakes
 - Znemožnění tvorby
 - Užití deepfakes
 - Detekce
 - Dokazování authenticity
 - Vzdělávání

Table 1: Connection of mitigation methods to the areas in deepfake lifecycle. The ● / ○ symbol represents whether a method applies to a specific area.

Method	Area		
	Preparation	Creation	Usage
Digital watermarking	●	●	●
Legal regulations	●	●	●
Obstructing deepfake creation	●	●	○
Deepfake detection	○	○	●
Forensics analysis	○	○	●
Proof of authenticity	○	○	●
Human education	○	○	●

- Omezení možností zneužít deepfakes
- Některé společnosti volí prevenci omezením služeb
 - Open AI:
 - "Phasing out voice-based authentication as a security measure for accessing bank accounts and other sensitive information"
- Vhodné nastavení procesů snižuje možnost provedení útoku
 - Stejně platí pro social engineering útoky

Nejrozšířenější ochrana - detekce



- EER - bod kde APCER = BPCER
 - APCER = Attack Presentation Classification Error
 - Počet útočných (deepfake) vzorek klasifikovaných jako bonafide
 - BPCER = Bona fide Presentation Classification Error
 - Počet bonafide vzorek klasifikovaných jako útok (deepfake)
- Jak důvěryhodná je 99% přesnost detekce?

Dataset	EER (%)
ASVspoof 2015	0.09
ASVspoof 2019 LA	0.06
ASVspoof 2021 LA	0.56
ASVspoof 2021 DF	1.16
WaveFake	1.30
In-the-wild	3.10

- *”Automatic Speaker Verification and Spoofing Countermeasures Challenge”*
- Nejznámější iniciativa zaměřená na vývoj metod pro detekci deepfake audia
- Každé 2 roky nová výzva
 - Sdružuje odborníky z oblasti
 - Zaměřená na vývoj nových (lepších) algoritmů pro detekci
 - ne systémů!
- Datasets z výzev jsou považovány za standard pro trénink a testování detektorů

- Akademické publikace produkují algoritmy (metody)
 - Specifický postup nebo sada výpočtů, která slouží k analýze zvukových dat a rozpoznání pravosti
 - Testován a optimalizován v laboratorním prostředí
 - Potřeba upravit pro použití v praxi
- My ale potřebujeme detekční systémy
 - kompletní implementace detekčního algoritmu v reálném světě
 - nejen samotný algoritmus, ale i další komponenty, jako je rozhraní pro příjem a zpracování zvukových dat
- Algoritmus -> systém
 - Technologické a legislativní překážky

9. Do not use our Services in any manner contrary to ElevenLabs' policies, purpose or mission.

This includes:

k) Using any part of our Services or their Output as input for any machine learning or training of artificial intelligence models.

l) Using any part of our Services or their Output as part of a dataset that may be used for training, fine-tuning, developing, testing, or improving any machine learning or artificial intelligence technology.

9. Do not use our Services in any manner contrary to ElevenLabs' policies, purpose or mission.

This includes:

- k) Using any part of our Services or their Output as input for any machine learning or training of artificial intelligence models.
- l) Using any part of our Services or their Output as part of a dataset that may be used for training, fine-tuning, developing, testing, or improving any machine learning or artificial intelligence technology.

Year	Only 2019 LA?	Frontend	EER (%) on 2019 LA
[1] 2022	Yes	waveform	0.64
[2] 2021	Yes	SincNet	0.83
[3] 2022	No (3DBs)	SincNet	0.93
[4] 2022	Yes	LFCC	0.98
[5] 2021	Yes	SincNet	1.06
[6] 2022	Yes	LFCC	1.06
[7] 2021	Yes	LFCC	1.07
[8] 2022	Yes	wav2vec 2.0	1.08
[9] 2022	No (4DBs)	wav2vec 2.0	1.28

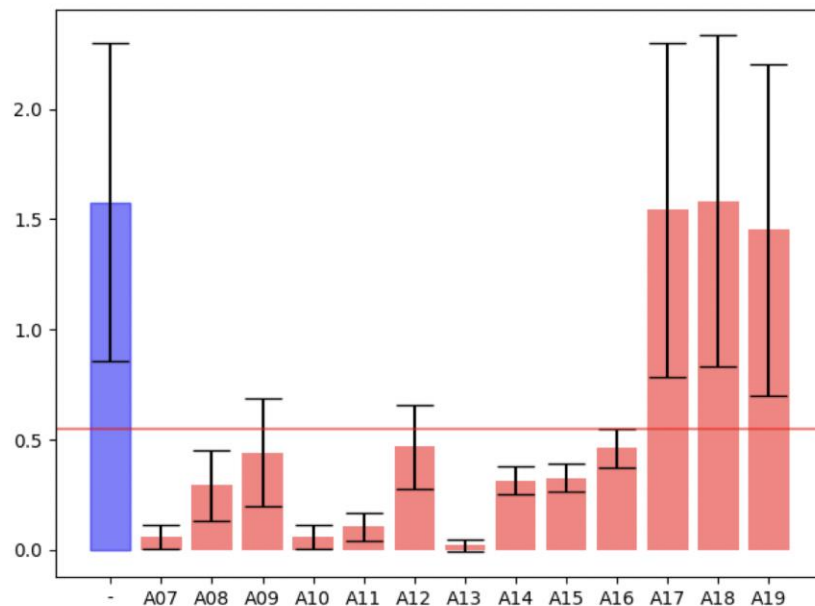
Year	Only 2019 LA?	Frontend	EER (%) on 2019 LA
[10] 2024	Yes	wav2vec 2.0	0.39
[11] 2024	Yes	wav2vec 2.0	0.51
[12] 2024	Yes	Masked autoencoder	0.25
[13] 2024	No (3 DBs)	wav2vec 2.0	0.17
[14] 2024	No (7 DBs)	wav2vec 2.0	0.40
[15] 2024	No (3 DBs)	wav2vec 2.0	0.07
[16] 2024	No (3 DBs)	wavLM	0.65
[17] 2024	No (3 DBs)	wavLM	0.23
[18] 2024	No (5 DBs)	wav2vec 2.0	2.36

Year	Only 2019 LA?	Frontend	EER (%) on 2019 LA
[1] 2022	Yes	waveform	0.64
[2] 2021	Yes	SincNet	0.83
[3] 2022	No (3DBs)	SincNet	0.93
[4] 2022	Yes	LFCC	0.98
[5] 2021	Yes	SincNet	1.06
[6] 2022	Yes	LFCC	1.06
[7] 2021	Yes	LFCC	1.07
[8] 2022	Yes	wav2vec 2.0	1.08
[9] 2022	No (4DBs)	wav2vec 2.0	1.28

Year	Only 2019 LA?	Frontend	EER (%) on 2019 LA
[10] 2024	Yes	wav2vec 2.0	0.39
[11] 2024	Yes	wav2vec 2.0	0.51
[12] 2024	Yes	Masked autoencoder	0.25
[13] 2024	No (3 DBs)	wav2vec 2.0	0.17
[14] 2024	No (7 DBs)	wav2vec 2.0	0.40
[15] 2024	No (3 DBs)	wav2vec 2.0	0.07
[16] 2024	No (3 DBs)	wavLM	0.65
[17] 2024	No (3 DBs)	wavLM	0.23
[18] 2024	No (5 DBs)	wav2vec 2.0	2.36

Jediný systém s řešením biasu ticha v nahrávkách

(c) This plot shows the average silence duration in seconds per audio file for the data split 'eval'.

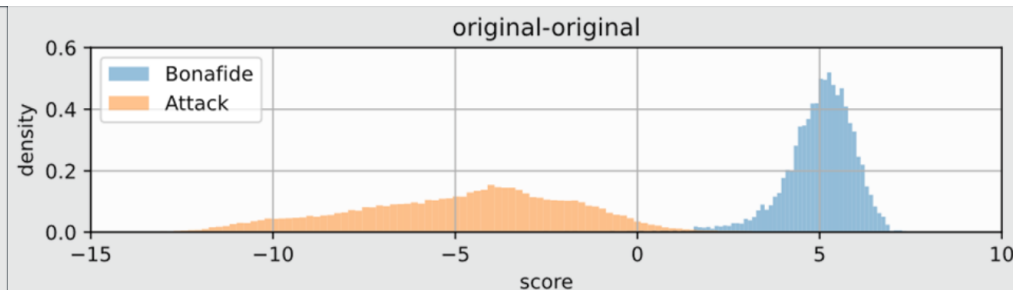
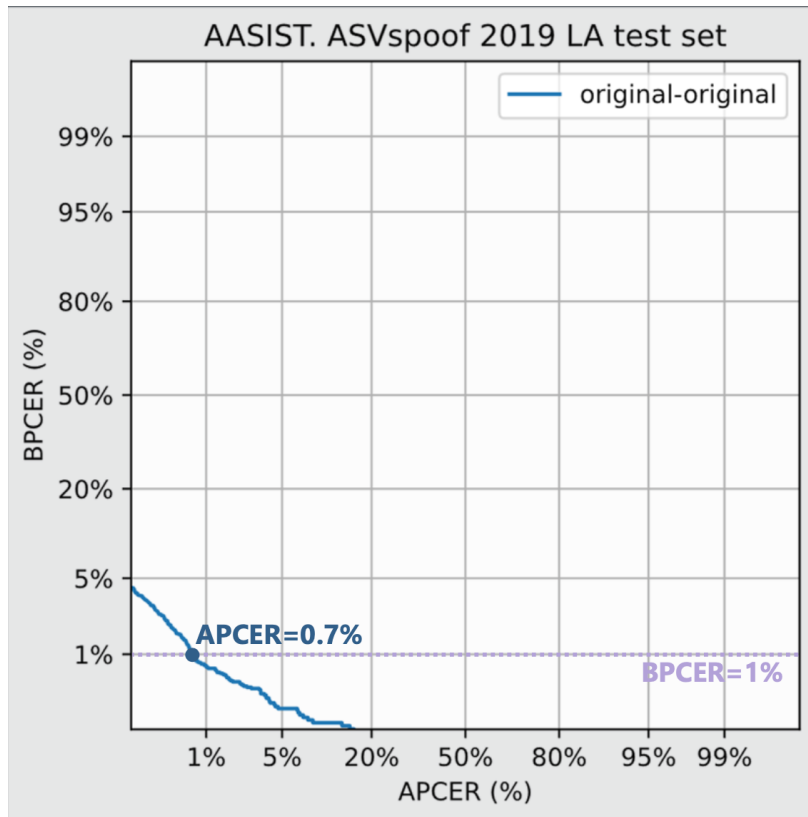


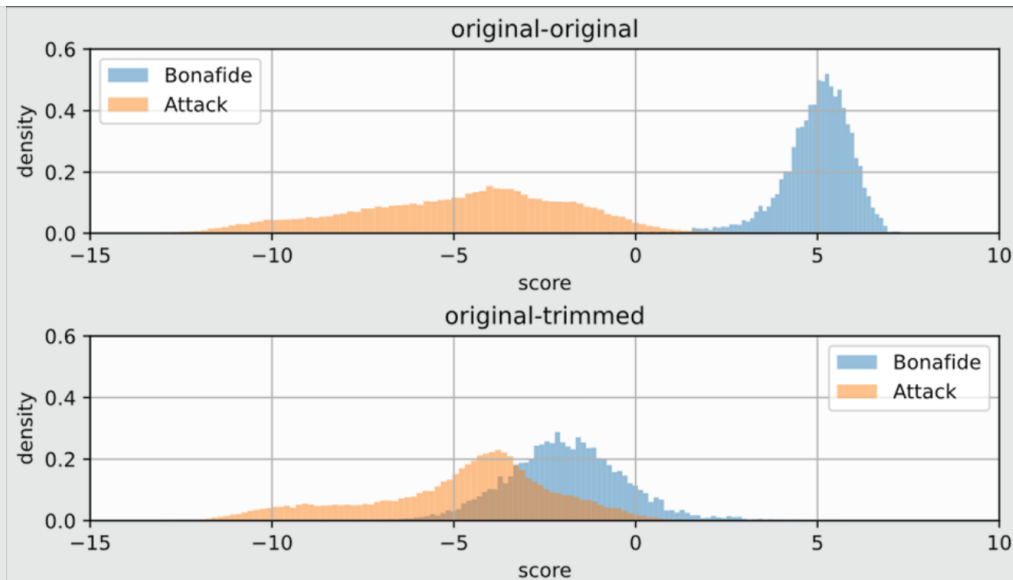
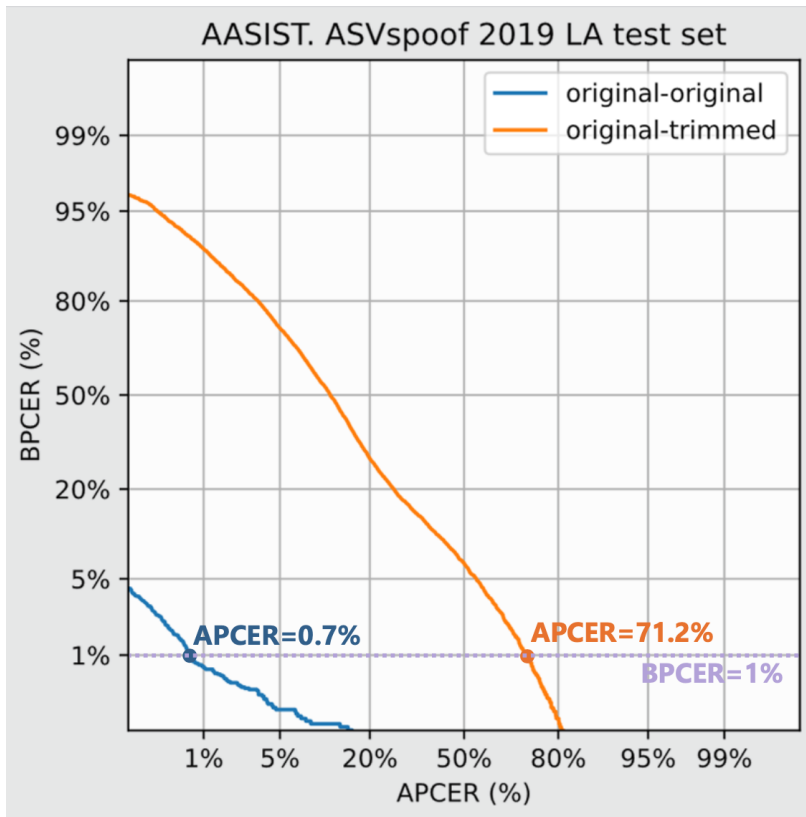
[1] Müller, N., Dieckmann, F., Czempin, P., Canals, R., Böttinger, K., Williams, J. (2021) Speech is Silver, Silence is Golden: What do ASVspoof-trained Models Really Learn? Proc. 2021 Edition of the Automatic Speaker Verification and Spoofing Countermeasures Challenge, 55-60, doi: 10.21437/ASVSPPOOF.2021-9

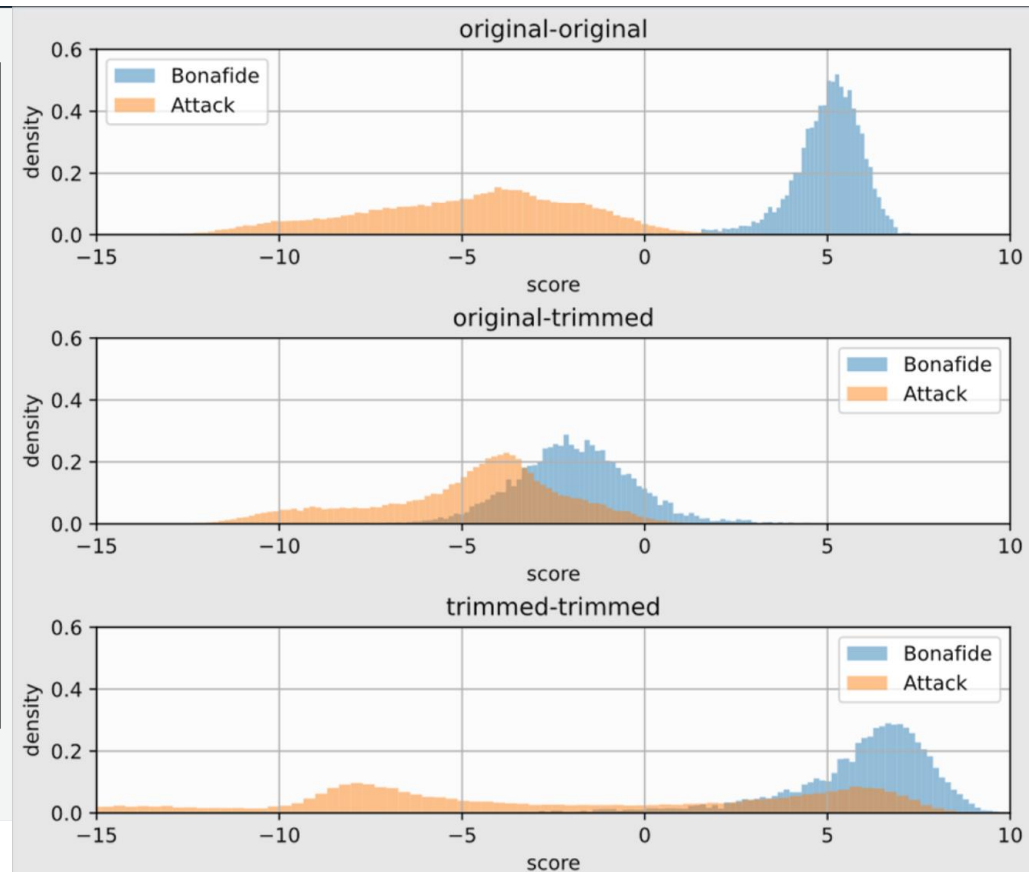
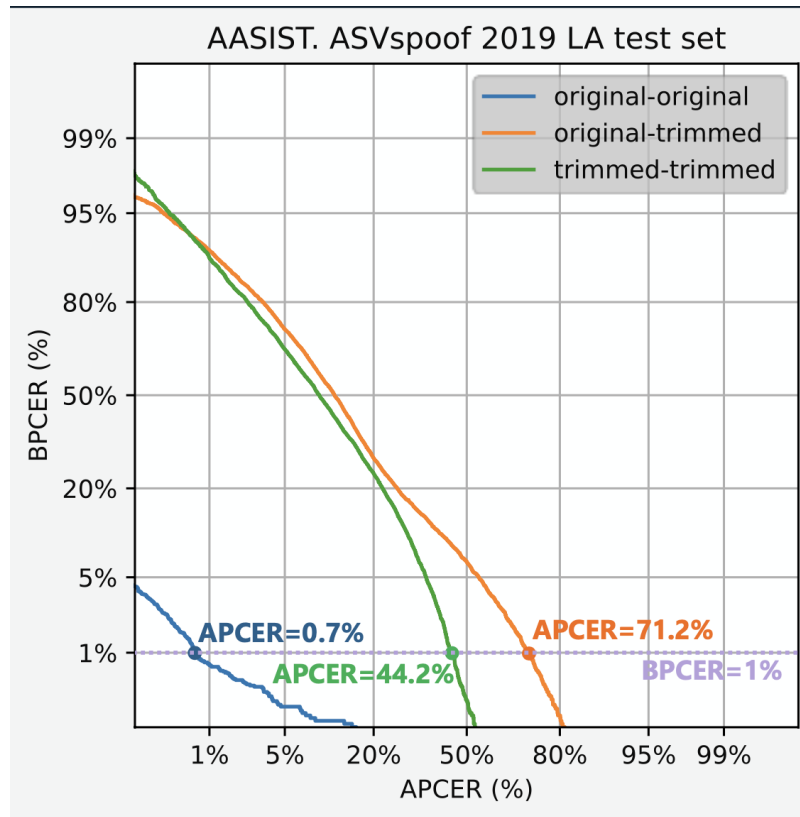
- B. Chettri et al., “A Deeper Look at Gaussian Mixture Model Based Anti-Spoofing Systems,” ICASSP, 2018
 - *“Our analysis shows how system performance can depend on a simple class-dependent cue in the dataset: initial silence frames of zeros appear in the genuine signals but missing in the spoofed version.”*
- N. M. Müller et al, “Speech is Silver, Silence is Golden: What do ASVspoof-trained Models Really Learn?,” ASVspoof workshop, 2021
 - *“Results show that models trained on only the duration of the leading silence perform suspiciously well, and achieve up to 85% accuracy and an equal error rate (EER) of 15.1%”*
 - *“...we observe that even for established antispoofing models such as RawNet2, removing silence during training leads to comparatively worse performance.”*
- R. Geirhos et al, “Shortcut Learning in Deep Neural Networks,” in Nature Machine Intelligence, 2023
 - *“Shortcuts are decision rules that perform well on standard benchmarks but fail to transfer to more challenging testing conditions, such as real-world scenarios”*

- AASIST: Audio Anti-Spoofing using Integrated Spectro-Temporal Graph Attention Networks
- SOTA v té době. EER=0,83 % na databázi ASVspoof 2019 LA
- Při trénování i inferenci systém používá počáteční 4s úsek vstupního zvuku, včetně počátečního ticha

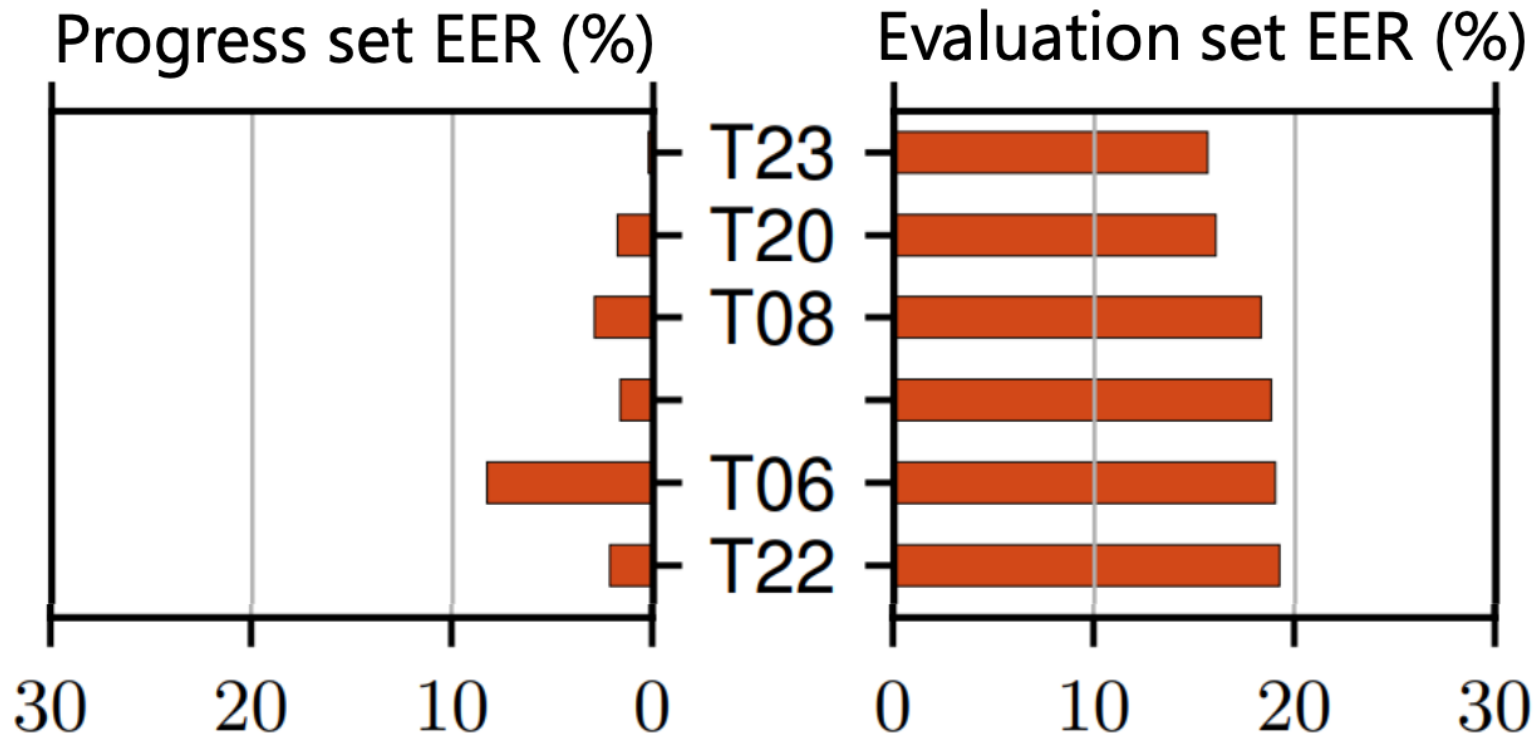
Condition	Trained on	Tested on
original-original	Original data	Original data
original-trimmed	Original data	Leading/trailing silences removed
trimmed-trimmed	Leading/trailing silences removed	Leading/trailing silences removed







Dataset	EER (%)	Unbiased EER (%)
ASVspoof 2015	0.09	-
ASVspoof 2019 LA	0.06	2.36
ASVspoof 2021 LA	0.56	9.47
ASVspoof 2021 DF	1.16	5.55
WaveFake	1.30	-
In-the-wild	3.10	-

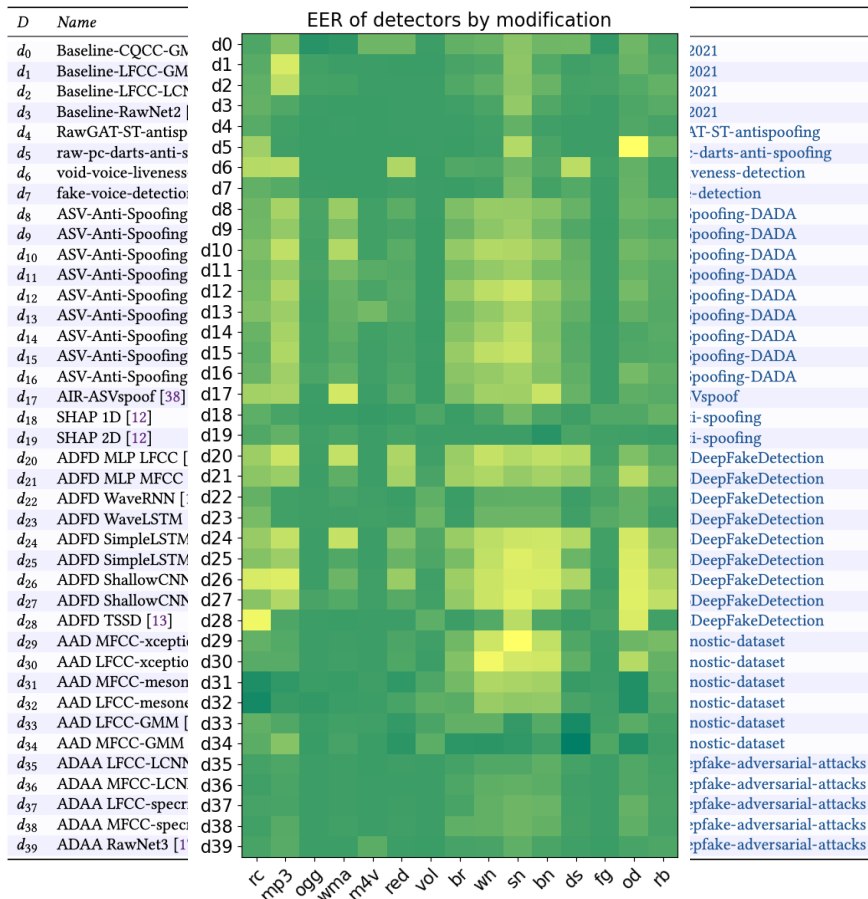


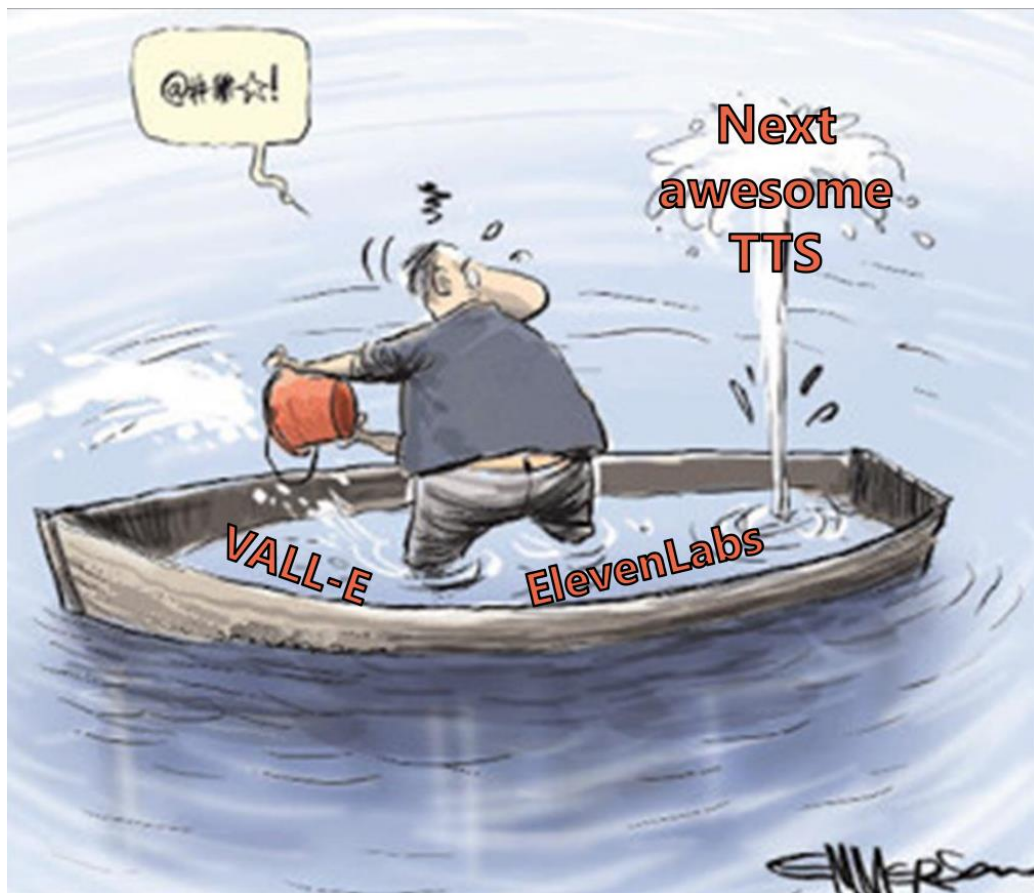
- Velké množství metod
- Slabá robustnost
- Těžká přenositelnost do praxe
- Jedno všeobecné řešení?

<i>D</i>	<i>Name</i>	<i>Link</i>
<i>d</i> ₀	Baseline-CQCC-GMM [6]	https://github.com/asvspoof-challenge/2021
<i>d</i> ₁	Baseline-LFCC-GMM [6]	https://github.com/asvspoof-challenge/2021
<i>d</i> ₂	Baseline-LFCC-LCNN [6]	https://github.com/asvspoof-challenge/2021
<i>d</i> ₃	Baseline-RawNet2 [6]	https://github.com/asvspoof-challenge/2021
<i>d</i> ₄	RawGAT-ST-antispoofing [15, 31]	https://github.com/eurecom-asp/RawGAT-ST-antispoofing
<i>d</i> ₅	raw-pc-darts-anti-spoofing [11]	https://github.com/eurecom-asp/raw-pc-darts-anti-spoofing
<i>d</i> ₆	void-voice-liveness-detection [1]	https://github.com/chislab/void-voice-liveness-detection
<i>d</i> ₇	fake-voice-detection [7]	https://github.com/dessa-oss/fake-voice-detection
<i>d</i> ₈	ASV-Anti-Spoofing-DADA LightCNN [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₉	ASV-Anti-Spoofing-DADA LightCNN DAT [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₀	ASV-Anti-Spoofing-DADA LightCNN DADA [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₁	ASV-Anti-Spoofing-DADA CGCNN [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₂	ASV-Anti-Spoofing-DADA CGCNN DAT [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₃	ASV-Anti-Spoofing-DADA CGCNN DADA [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₄	ASV-Anti-Spoofing-DADA ResNet [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₅	ASV-Anti-Spoofing-DADA ResNet DAT [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₆	ASV-Anti-Spoofing-DADA ResNet DADA [32]	https://github.com/Jijiang/ASV-Anti-Spoofing-DADA
<i>d</i> ₁₇	AIR-ASVspoof [38]	https://github.com/yzyouzhang/AIR-ASVspoof
<i>d</i> ₁₈	SHAP 1D [12]	https://github.com/gewanying/shap-anti-spoofing
<i>d</i> ₁₉	SHAP 2D [12]	https://github.com/gewanying/shap-anti-spoofing
<i>d</i> ₂₀	ADFD MLP LFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₁	ADFD MLP MFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₂	ADFD WaveRNN [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₃	ADFD WaveLSTM [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₄	ADFD SimpleLSTM LFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₅	ADFD SimpleLSTM MFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₆	ADFD ShallowCNN LFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₇	ADFD ShallowCNN MFCC [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₈	ADFD TSSD [13]	https://github.com/MarkHershey/AudioDeepFakeDetection
<i>d</i> ₂₉	AAD MFCC-xception [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₀	AAD LFCC-xception [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₁	AAD MFCC-mesonet-inception (spectrogram) [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₂	AAD LFCC-mesonet-inception (spectrogram) [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₃	AAD LFCC-GMM [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₄	AAD MFCC-GMM [16]	https://github.com/piotrkawa/attack-agnostic-dataset
<i>d</i> ₃₅	ADAA LFCC-LCNN [17]	https://github.com/piotrkawa/audio-deepfake-adversarial-attacks
<i>d</i> ₃₆	ADAA MFCC-LCNN [17]	https://github.com/piotrkawa/audio-deepfake-adversarial-attacks
<i>d</i> ₃₇	ADAA LFCC-specnet [17]	https://github.com/piotrkawa/audio-deepfake-adversarial-attacks
<i>d</i> ₃₈	ADAA MFCC-specnet [17]	https://github.com/piotrkawa/audio-deepfake-adversarial-attacks
<i>d</i> ₃₉	ADAA RawNet3 [17]	https://github.com/piotrkawa/audio-deepfake-adversarial-attacks

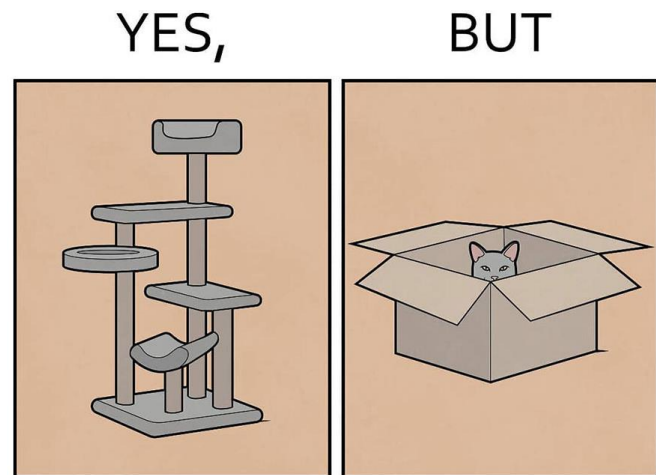
Table 2: Detection methods

- Velké množství metod
- Slabá robustnost
- Těžká přenositelnost do praxe
- Jedno všeobecné řešení?





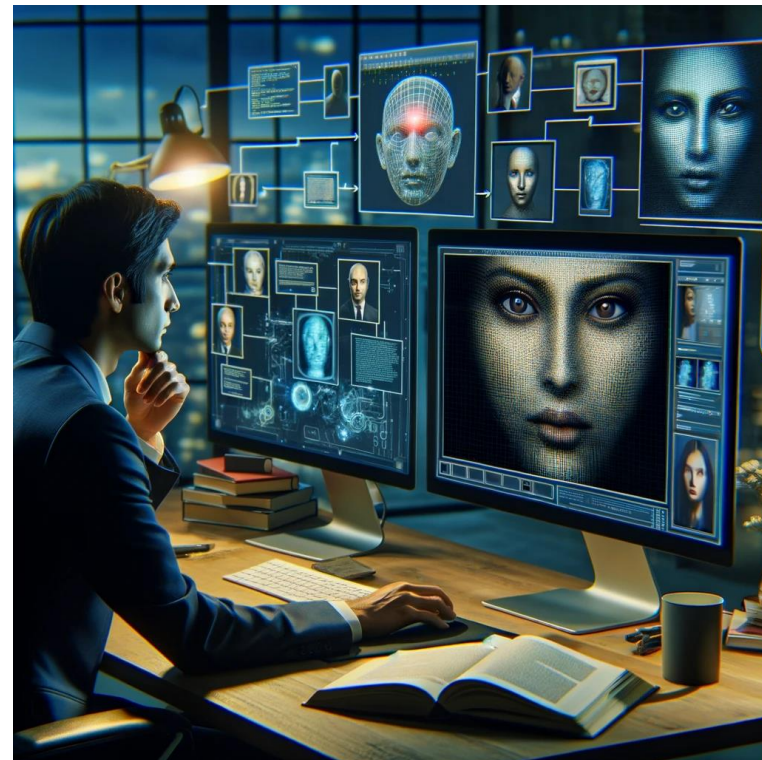
- Generalizace (a zkratky)
 - Nové typy útoků (nástrojů)
 - Variabilita dat
- Vysvětlitelnost
 - Proč?
- Certifikace
 - Jak moc věříme výsledkům?



<https://www.boredpanda.com/yes-but-comics-society-criticism-anton-gudim/>

© _yes_but

- Lidi nedokážou rozpoznat pravé a deepfake média
- Možnost zmírnění hrozeb
 - Zvyšování povědomí o technologii deepfakes
 - Opakovaná expozice deepfake médiím zvyšuje rozeznávací schopnost
 - Možnost inspirace cvičnými phishing kampaněmi



- Deepfakes představují hrozbu pro lidi a biometrické systémy
- Lidi nedokážou rozlišit deepfakes od pravých nahrávek
- Biometrické systémy nemají přirozenou ochranu před deepfakes
- Tvorba deepfakes je čím dál jednodušší a dostupnější
- Potřeba věnovat se kvalitním metodám pro detekci



Kamil Malinka, Anton Firc
malinka@fit.vut.cz, ifirc@fit.vut.cz



Security@FIT
<https://www.fit.vut.cz/research/group/security@fit/.en>

- Proč se nebát AI
 - AI pohledem nejen českých odborníků
 - Kapitola o bezpečnostních dopadech deepfakes
 - <https://www.jota.cz/proc-se-nebat-umele-intelligence.html>
- Česká Asociace Umělé Inteligence
 - Obranná strategie pro české firmy
 - <https://asociace.ai/wp-content/uploads/2023/12/DEEPFAKE-2024-CAUI.pdf>

