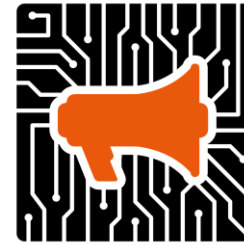




ELTE | PPK  
PEDAGÓGIAI ÉS PSZICHOLÓGIAI KAR



DEZINFORMÁCIÓ  
ÉS MESTERSÉGES  
INTELLIGENCIA  
KUTATÓCSOPORT

# Mesterséges intelligencia és dezinformáció – utópiák és disztópiák

Krekó Péter, PhD.

Habilitált egyetemi docens, ELTE PPK Szociálpszichológia Tanszék/  
Dezinformáció és Mesterséges Intelligencia Kutatócsoport

„Vissza a jövőbe!”- A közösségi média és az AI hatása a történelemre  
Történelemtanárok Egylete, 2024. október 5.

# Fő kérdések

---

- Hogyan hat a dezinformáció előállításával kapcsolatos technológiai fejlődés az információ kognitív feldolgozására, az érzelmi reakciókra?
- Milyen releváns pszichológiai vizsgálatok vannak a területen?
- Milyen gyakorlati jelentősége van mindennek?
- Milyen irányba fordíthatja ez a történelmet?

# Miről lesz szó?

---

1) Előállítás

2) Terjesztés

3) Detekció/Prevenció

+ Következtetések

People: fearing AI takeover  
AI:

| Clipboard |     | Font                                                                                                                                                                                                                                                       |         | Alignm |  |
|-----------|-----|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|--------|--|
| B2        |     |    | Febuary |        |  |
|           | A   | B                                                                                                                                                                                                                                                          | C       |        |  |
| 1         | JAN | January                                                                                                                                                                                                                                                    |         |        |  |
| 2         | FEB | Febuary                                                                                                                                                                                                                                                    |         |        |  |
| 3         | MAR | Maruary                                                                                                                                                                                                                                                    |         |        |  |
| 4         | APR | Apruary                                                                                                                                                                                                                                                    |         |        |  |
| 5         | MAY | Mayuary                                                                                                                                                                                                                                                    |         |        |  |
| 6         | JUN | Junuary                                                                                                                                                                                                                                                    |         |        |  |

**‘Sapiens’ author says AI is an alien threat that could wipe us out: ‘Instead of coming from outer space, it’s coming from California’**











ELTE | PPK  
PEDAGÓGIAI ÉS PSZICHOLÓGIAI KAR



DEZINFORMÁCIÓ  
ÉS MESTERSÉGES  
INTELLIGENCIA  
KUTATÓCSOPORT

# 1) A tartalom előállítása

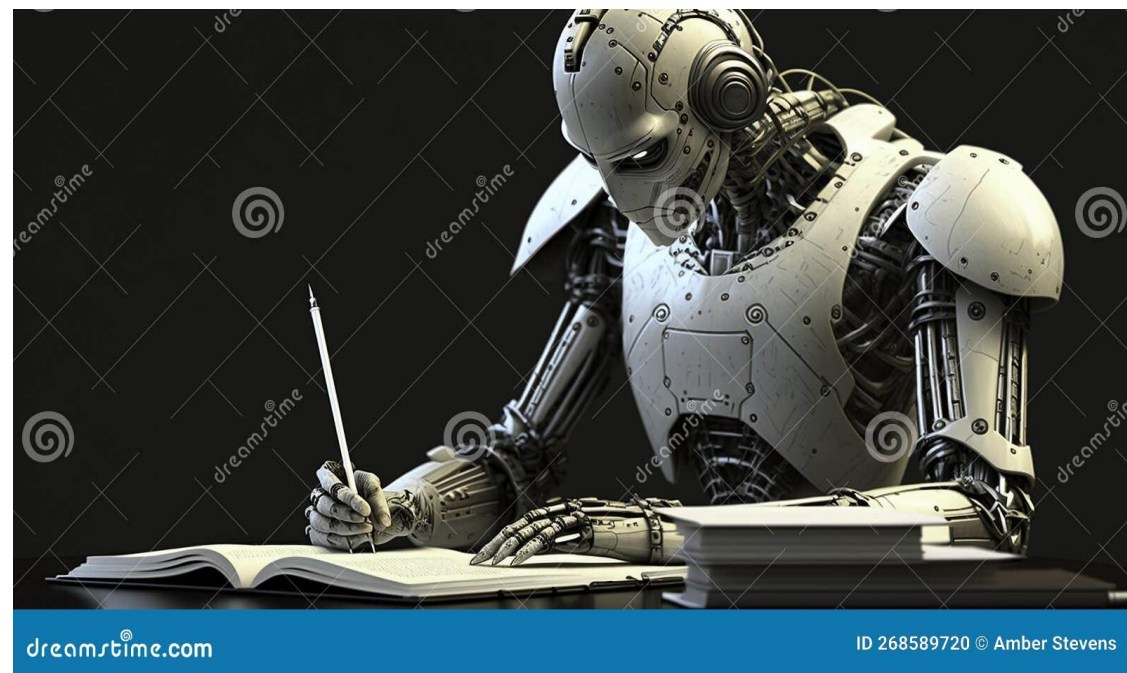
# A generatív AI sokszor meggyőzőbben csap be minket

Hitelesebbnek fogadjuk el a GPT-3, mint az emberek által írt (dez) információt (Spitale et al, 2023)

A GPT-3 által írt twitter-üzenetek az emberek számára nem megkülönböztethetőek a személyek által írtaktól (u.o.)

A generatív AI által létrehozott üzenetek hatékonyan meggyőzik a befogadókat sokféle politikai kérdésben (Bai et al., preprint).

Az AI által generált üzeneteket ténszerűbbnek, logikusabbnak, DE kevésbé dühösnek, egyedinek, és kevésbé „sztorizósnak” ítélték a kérdezettek (u.o.).





# A bullshit-problematika (misinformation)

## ON BULLSHIT

Harry G. Frankfurt

## Why chatbots are bound to spout bullshit

Some of what they say is true, but only as a byproduct of learning to seem believable



© Guillem Casaus

## Who falls for fake news? The roles of bullshit receptivity, overclaiming, familiarity, and analytic thinking

Gordon Pennycook<sup>1</sup>  | David G. Rand<sup>2,3</sup>

<sup>1</sup>Hill/Levene Schools of Business, University of Regina, Regina, Canada

<sup>2</sup>Sloan School, Massachusetts Institute of Technology, Cambridge, Massachusetts

<sup>3</sup>Department of Brain and Cognitive Sciences, Massachusetts Institute of Technology, Cambridge, Massachusetts

### Correspondence

Gordon Pennycook, Hill/Levene Schools of Business, University of Regina, 3737 Wascana Pky, Regina, Saskatchewan, Canada, S4S 0A2  
Email: gordon.pennycook@uregina.ca

### Funding information

Social Sciences and Humanities Research Council of Canada; Miami Foundation

### Abstract

**Objective:** Fake news represents a particularly egregious and direct avenue by which inaccurate beliefs have been propagated via social media. We investigate the psychological profile of individuals who fall prey to fake news.

**Method:** We recruited 1,606 participants from Amazon's Mechanical Turk for three online surveys.

**Results:** The tendency to ascribe profundity to randomly generated sentences—pseudo-profound bullshit receptivity—correlates positively with perceptions of fake news accuracy, and negatively with the ability to differentiate between fake and real news (media truth discernment). Relatedly, individuals who overclaim their level of knowledge also judge fake news to be more accurate. We also extend previous research indicating that analytic thinking correlates negatively with perceived accuracy by showing that this relationship is not moderated by the presence/absence of the headline's source (which has no effect on accuracy), or by familiarity with the headlines (which correlates positively with perceived accuracy of fake and real news).

**Conclusion:** Our results suggest that belief in fake news may be driven, to some extent, by a general tendency to be overly accepting of weak claims. This tendency, which we refer to as reflexive open-mindedness, may be partly responsible for the prevalence of epistemically suspect beliefs writ large.

### KEYWORDS

analytic thinking, bullshit receptivity, fake news, news media, social media

# Pseudo-profound bullshit skála

---

Rezgünk, rezgünk, újjászületünk. Elektromágneses rezonanciaként létezünk.

Az egymásra taltságot elmúló cselekedetekben gyökerezik.

Az örök nyugalom újjászületik a végtelen emberi megfigyelésben.

A végtelen a karmikus valóság tükröződése.

A rejtett jelentés a csodák szimbolikus ábrázolásától függ.

A transzcendencia átalakítja a kiválóság kapuját.

Egy átalakulás érző ébredésének közepén vagyunk, amely összehangol minket magával a bioszférával.

A hálózat közeledik a fordulóponthoz. A kozmikus tudás ébredése lesz a jövő.

Nem engedhetjük meg többé magunknak, hogy megszakítással éljünk.

A megszakítás a belátás ellentéte.

A tudat a kvantumenergia morfikus rezonanciájából áll.

A kvantum az erő ébredését jelenti



uilt entirely with AI imagery







Home » Research

## Types, sources, and claims of COVID-19 misinformation



A man wearing a face mask uses his phone in the Times Square subway station. REUTERS/Lucas Jackson

Dr. J. Scott Brennen · Felix Simon · Dr. Philip N. Howard · Prof. Rasmus Kleis Nielsen  
Tuesday 7 April 2020



Russia in South Africa  
@EmbassyofRussia

Today is 22nd anniversary of the day when NATO brought death & destruction to #Yugoslavia by attacking a sovereign state in total disregard for intl law. NATO's 78 days long "humanitarian intervention" took lives of 2000 innocent people, turned thousands more into refugees

[Bejegyzés lefordítása](#)



MFA Russia és 8 másik felhasználó

# Israel-Hamas war: Beware of the cheap fake videos on the Internet

Two misinformation experts explain why and how you can develop the power to resist these deceptions.

Sam Wineburg and Michael Caulfield



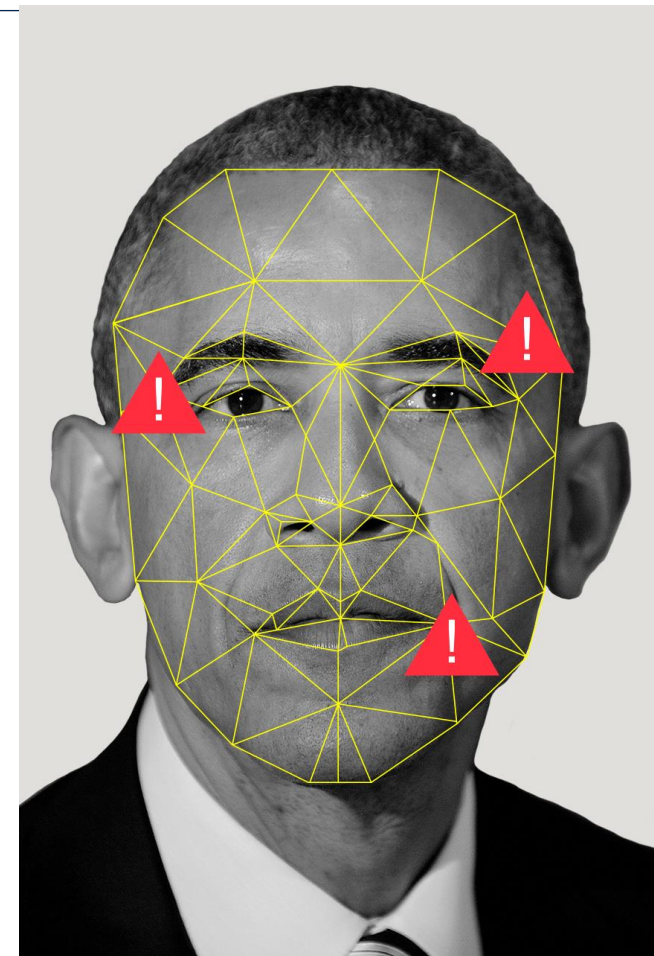
A makeshift operating theater area inside Al-Shifa hospital during the Israeli ground operation around the hospital, in Gaza City, on Nov 12. PHOTO: REUTERS

UPDATED 2023. NOV. 18. DE. 6:04 SGT -



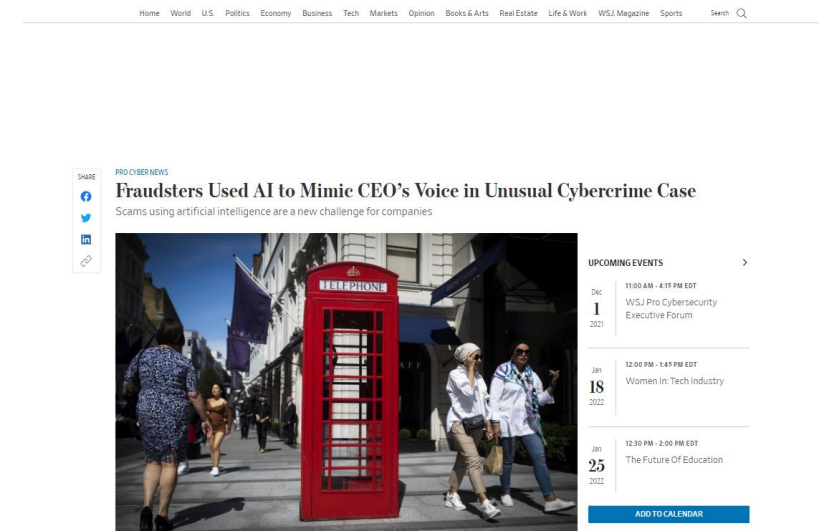
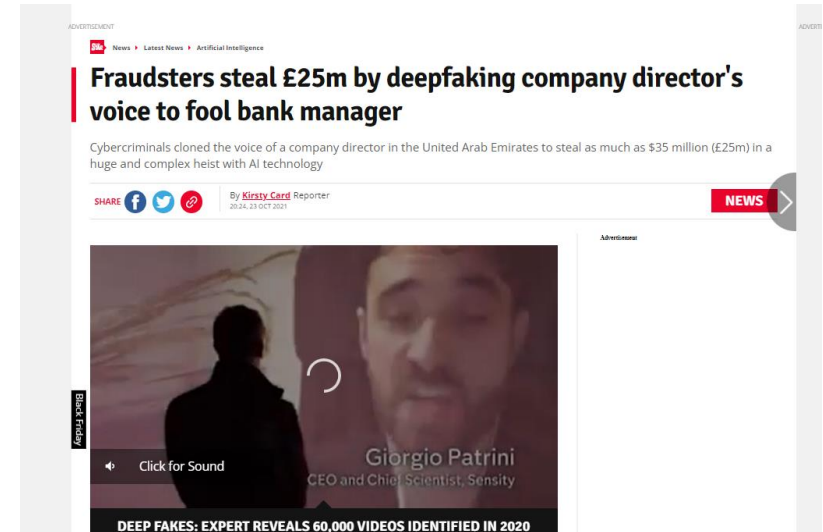
# Deepfake-ek

- Az álhírek „nehézfegyverzete” (Veszelszki, 2020)
- A “deep learning” és “fake” szavak kombinációja (+deep voice)
- Ultrarealisztikus, gépi tanulással, azon belül mély neurális hálózatok felhasználásával létrehozott hamis audiovizuális tartalom
- Valós képi és hangzó anyag felhasználásával készül
- Sokszor valós háttéren, csak minimális módosítással
- **Erősebb érzelmi reakció kiváltására alkalmas** (Arsenault, 2020)
- Az emberek hajlamosak silány minőségű valótlan tartalmakat is elhinni, ha azt saját pártjuktól származik (Whyte, 2020).
- DE analitikus képességek és a politikai érdeklődés segít a felismerésükben (Appel, Prietzel, 2023)



# A deepfake veszélyei: a valóra vált post-truth disztópia?

- 2030-ra [várhatóan](#) szabad szemmel detektálhatatlanná, és mindenki számára előállíthatóvá válik ez a technológia
- Az emberek nem tudják megkülönböztetni a valóságtól: „csak a szememnek hiszek”
- Hatékonyabbak a hagyományos fake-eknél, akkor is hat, ha tudjuk hogy hamis (!)
- Alkalmazásai:
  - Személyes bosszú/cyberbullying; Személyes erőszak: pl. Whatsapp-lincselések Indiában
  - Politikai: választások befolyásolása, diplomáciai konfliktusok és fegyveres konfliktusok szítása,
  - Jogi: megtévesztő marketingkampányok, bizonyíték jogi eljárásban
  - Gazdasági: megtévesztő marketingkampányok, csalások, vállalatok elleni fake news kampányok





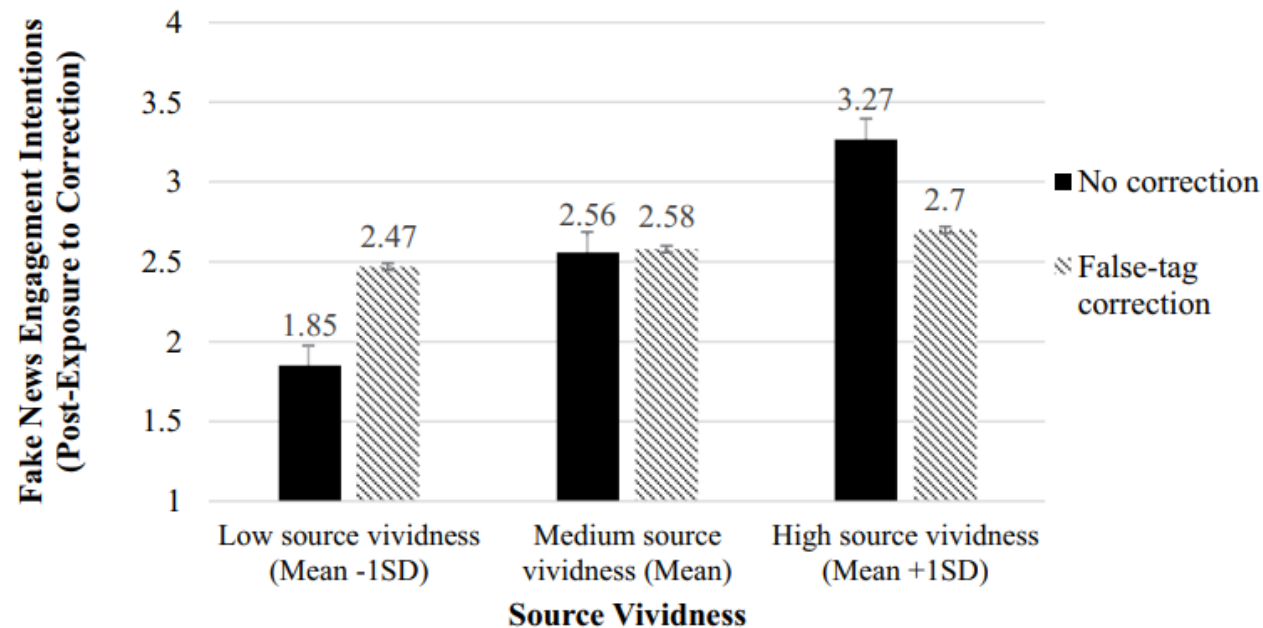
# A Deepfake és az episztemikus bizonytalanság

A jelenség a megvezetés új lehetőségeit kínálja: a hamis információt állítja be igaznak vagy pedig a hiteles információt hamisnak állíthatjuk be (Bontridder & Poulet, 2021).

A deepfake-ek veszélyével való riogatás nem javítja a valósi és manipulált videók elkülönítésének képességét, csak a szkepszist erősíti (Ternovski et al., 2021).

To cite this article: Jiyoung Lee & Soo Yun Shin (2021): Something that They Never Said: Multimodal Disinformation and Source Vividness in Understanding the Power of AI-Enabled Deepfake News, *Media Psychology*, DOI: [10.1080/15213269.2021.2007489](https://doi.org/10.1080/15213269.2021.2007489)

To link to this article: <https://doi.org/10.1080/15213269.2021.2007489>



**figure 1.** Interaction between source vividness and false-tag correction (vs. no correction) on fake news engagement intentions (study 2; abortion). Error bars are 95% confidence intervals.



ELTE | PPK  
PEDAGÓGIAI ÉS PSZICHOLÓGIAI KAR



DEZINFORMÁCIÓ  
ÉS MESTERSÉGES  
INTELLIGENCIA  
KUTATÓCSOPORT

## 2) Tartalom terjesztése

# A dezinformáció terjesztése

---

Demográfiai, és pszichológiai profilírozás (pl. Cambridge Analytica; Youyou, 2015).

Szegmentálás és mikrotargetálás (pl. (Dobber et al., 2020). hatékonyabb visszhangkamrába zárás, intellektuális izoláció, naív szociológia,

Szociális botok:

- 2016-os amerikai elnökválasztási kampány utolsó hónapjában a Twitteren megjelenő posztok egyötödét botok generálták és terjesztették el (Masood és mtsai., 2023; Sebastian, 2023).
- A hiteltelen források hatásának növelésében rendkívül hatásosak lehetnek a social bot-ok, az emberek pedig gyakran nem azonosítják őket, és megosztják/vitába szállnak velük (Shao et al., 2018).



### 3) Az álhírek detekciója/azonosítása/cáfolata



ELTE | PPK  
PEDAGÓGIAI ÉS PSZICHOLÓGIAI KAR



DEZINFORMÁCIÓ  
ÉS MESTERSÉGES  
INTELLIGENCIA  
KUTATÓCSOPORT

### 3) Az álhírek detekciója/cáfolata

# Stratégiák a dezinformáció elleni fellépésben (Krekó, 2020)

|                                      | Dezinformációval való találkozás előtt (Prevenció)                                            | Dezinformációval való találkozás után (utólagos kezelés)                 |
|--------------------------------------|-----------------------------------------------------------------------------------------------|--------------------------------------------------------------------------|
| Kínálati oldalt célzó beavatkozások  | <b>Megelőző csapás</b><br>(pl. eltávolítás)                                                   | <b>Válaszcsapás</b><br>(pl. Álhírek címkézése, a forrás hiteltelenítése) |
| Keresleti oldalt célzó beavatkozások | <b>Immunizálás</b><br>(pl. előzetes figyelmeztetés, beoltás, az analitikus készség erősítése) | <b>Gyógyítás</b><br>(„Kognitív infiltráció”<br>Cáfolat/fact-checking)    |



# Leveraging ChatGPT for Efficient Fact-Checking

Emma Hoes<sup>‡</sup>, Sacha Altay<sup>†</sup> and Juan Bermeo<sup>§</sup>

May 29, 2023

## Abstract

We explore the potential of automated online content moderation through Large Language Models (LLMs) as a remedy to the overwhelming surge of online content that fact-checking organizations cannot verify. Although LLMs, such as ChatGPT, may exacerbate this problem by facilitating content production, they could also be leveraged to enhance the efficiency and expediency of fact-checking processes. We conducted a systematic evaluation to measure ChatGPT's fact-checking performance by submitting 21,152 fact-checked statements to ChatGPT as a zero-shot classification. We find that ChatGPT is able to accurately categorize statements as true or false in 69% of cases. Addressing memorization, the performance of ChatGPT is similar on claims that have not been fact-checked or are post its training data cutoff-date. These findings demonstrate the potential of ChatGPT to help label misinformation and deepen our comprehension of how LLMs could improve content moderation practices, complementing the crucial work of human fact-checking experts in upholding the accuracy of information dissemination.

**Keywords**— Misinformation, Fact-checking, Artificial Intelligence, AI, Fake News, ChatGPT, LLMs, Content Moderation, PolitiFact

**Word count:** 3353

# Pseudoscience Detection Using a Pre-trained Transformer Model with Intelligent ReLabeling

Thilina C. Rajapakse and Ruwan D. Nawarathna

**Abstract** Often dismissed as a harmless pastime for the gullible, pseudoscience nonetheless has devastating effects, including the loss of life. Its menace is amplified by the difficulty of differentiating pseudoscience from science, especially for the untrained eye. A novel method recognizes pseudoscience in the text by utilizing a fine-tuned Robustly Optimized Bidirectional Encoder Representation from Transformers Approach (RoBERTa) model. The dataset of 112,720 full-text articles used in this work is made publicly available to remedy the lack of datasets related to pseudoscience. A novel technique, Intelligent ReLabeling (IRL), is employed to minimize mislabeled data, enabling the rapid creation of high-quality textual datasets. IRL eliminates the need for expensive manual verification processes and minimizes domain expertise requirements in many applications. The final model trained with IRL achieves an  $F_1$  score of 0.929 on a separate manually labeled test dataset.

**Keywords** Natural language processing · Classification · Data cleaning · Pseudoscience dataset



TABLE II  
SUMMARY OF PROMINENT DEEPPAKE DETECTION METHODS

| Methods                                       | Classifiers/<br>Techniques                         | Key Features                                                                                                                                                                                                                                                                                                                                                           | Dealing<br>with   | Datasets Used                                                                                                                                                                                                                                  |
|-----------------------------------------------|----------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Eye blinking [99]                             | LRCN                                               | <ul style="list-style-type: none"> <li>- Use LRCN to learn the temporal patterns of eye blinking.</li> <li>- Based on the observation that blinking frequency of deepfakes is much smaller than normal.</li> </ul>                                                                                                                                                     | Videos            | Consist of 49 interview and presentation videos, and their corresponding generated deepfakes.                                                                                                                                                  |
| Intra-frame and temporal inconsistencies [98] | CNN and LSTM                                       | CNN is employed to extract frame-level features, which are distributed to LSTM to construct sequence descriptor useful for classification.                                                                                                                                                                                                                             | Videos            | A collection of 600 videos obtained from multiple websites.                                                                                                                                                                                    |
| Using face warping artifacts [103]            | VGG16 [101]<br>ResNet50, 101 or 152 [102]          | Artifacts are discovered using CNN models based on resolution inconsistency between the warped face area and the surrounding context.                                                                                                                                                                                                                                  | Videos            | <ul style="list-style-type: none"> <li>- UADFV [104], containing 49 real videos and 49 fake videos with 32752 frames in total.</li> <li>- DeepfakeTIMIT [64]</li> </ul>                                                                        |
| MesoNet [94]                                  | CNN                                                | <ul style="list-style-type: none"> <li>- Two deep networks, i.e. Meso-4 and MesoInception-4 are introduced to examine deepfake videos at the mesoscopic analysis level.</li> <li>- Accuracy obtained on deepfake and FaceForensics datasets are 98% and 95% respectively.</li> </ul>                                                                                   | Videos            | Two datasets: deepfake one constituted from online videos and the FaceForensics one created by the Face2Face approach [111].                                                                                                                   |
| Capsule-forensics [106]                       | Capsule networks                                   | <ul style="list-style-type: none"> <li>- Latent features extracted by VGG-19 network [101] are fed into the capsule network for classification.</li> <li>- A dynamic routing algorithm [108] is used to route the outputs of three convolutional capsules to two output capsules, one for fake and another for real images, through a number of iterations.</li> </ul> | Videos/<br>Images | Four datasets: the Idiap Research Institute replay-attack [109], deepfake face swapping by [94], facial reenactment FaceForensics [110], and fully computer-generated image set using [112].                                                   |
| Head poses [104]                              | SVM                                                | <ul style="list-style-type: none"> <li>- Features are extracted using 68 landmarks of the face region.</li> <li>- Use SVM to classify using the extracted features.</li> </ul>                                                                                                                                                                                         | Videos/<br>Images | <ul style="list-style-type: none"> <li>- UADFV consists of 49 deep fake videos and their respective real videos.</li> <li>- 241 real images and 252 deep fake images from DARPA MediFor GAN Image/Video Challenge.</li> </ul>                  |
| Eye, teach and facial texture [114]           | Logistic regression and neural network             | <ul style="list-style-type: none"> <li>- Exploit facial texture differences, and missing reflections and details in eye and teeth areas of deepfakes.</li> <li>- Logistic regression and neural network are used for classifying.</li> </ul>                                                                                                                           | Videos            | A video dataset downloaded from YouTube.                                                                                                                                                                                                       |
| Spatio-temporal features with RCN [95]        | RCN                                                | Temporal discrepancies across frames are explored using RCN that integrates convolutional network DenseNet [79] and the gated recurrent unit cells [96]                                                                                                                                                                                                                | Videos            | FaceForensics++ dataset, including 1,000 videos [97].                                                                                                                                                                                          |
| Spatio-temporal features with LSTM [140]      | Convolutional bidirectional recurrent LSTM network | <ul style="list-style-type: none"> <li>- An XceptionNet CNN is used for facial feature extraction while audio embeddings are obtained by stacking multiple convolution modules.</li> <li>- Two loss functions, i.e. cross-entropy and Kullback-Leibler divergence, are used.</li> </ul>                                                                                | Videos            | FaceForensics++ [97] and Celeb-DF (5,639 deepfake videos) [141] datasets and the ASVSpooof 2019 Logical Access audio dataset [142].                                                                                                            |
| Analysis of PRNU [115]                        | PRNU                                               | <ul style="list-style-type: none"> <li>- Analysis of noise patterns of light sensitive sensors of digital cameras due to their factory defects.</li> <li>- Explore the differences of PRNU patterns between the authentic and deepfake videos because face swapping is believed to alter the local PRNU patterns.</li> </ul>                                           | Videos            | Created by the authors, including 10 authentic and 16 deepfake videos using DeepFaceLab [34].                                                                                                                                                  |
| Phoneme-viseme mismatches [133]               | CNN                                                | <ul style="list-style-type: none"> <li>- Exploit the mismatches between the dynamics of the mouth shape, i.e. visemes, with a spoken phoneme.</li> <li>- Focus on sounds associated with the M, B and P phonemes as they require complete mouth closure while deepfakes often incorrectly synthesize it.</li> </ul>                                                    | Videos            | Four in-the-wild lip-sync deepfakes from Instagram and YouTube (www.instagram.com/bill_posters_uk and youtube.be/VWMEDacz3L4) and others are created using synthesis techniques, i.e. Audio-to-Video (A2V) [63] and Text-to-Video (T2V) [134]. |

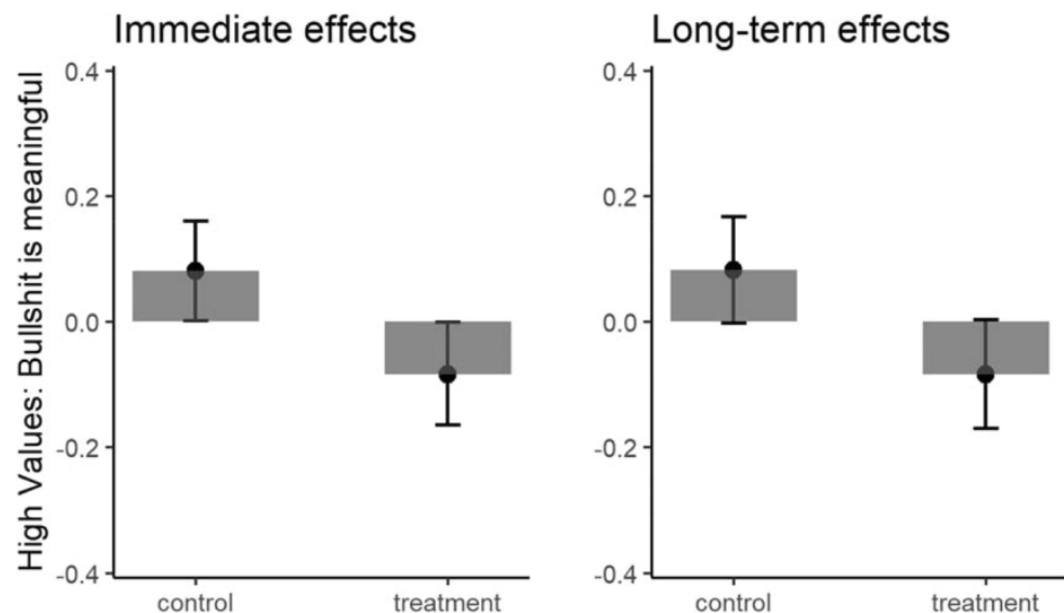
TABLE I  
SUMMARY OF NOTABLE DEEPPFAKE TOOLS

| Tools                             | Links                                                                                                                             | Key Features                                                                                                                                                                                                                                            |
|-----------------------------------|-----------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Faceswap                          | <a href="https://github.com/deepfakes/faceswap">https://github.com/deepfakes/faceswap</a>                                         | <ul style="list-style-type: none"> <li>- Using two encoder-decoder pairs.</li> <li>- Parameters of the encoder are shared.</li> </ul>                                                                                                                   |
| Faceswap-GAN                      | <a href="https://github.com/shaoanlu/faceswap-GAN">https://github.com/shaoanlu/faceswap-GAN</a>                                   | Adversarial loss and perceptual loss (VGGface) are added to an auto-encoder architecture.                                                                                                                                                               |
| Few-Shot Face Translation         | <a href="https://github.com/shaoanlu/fewshot-face-translation-GAN">https://github.com/shaoanlu/fewshot-face-translation-GAN</a>   | <ul style="list-style-type: none"> <li>- Use a pre-trained face recognition model to extract latent embeddings for GAN processing.</li> <li>- Incorporate semantic priors obtained by modules from FUNIT [42] and SPADE [43].</li> </ul>                |
| DeepFaceLab                       | <a href="https://github.com/iperov/DeepFaceLab">https://github.com/iperov/DeepFaceLab</a>                                         | <ul style="list-style-type: none"> <li>- Expand from the Faceswap method with new models, e.g. H64, H128, LIAEF128, SAE [44].</li> <li>- Support multiple face extraction modes, e.g. S3FD, MTCNN, dlib, or manual [44].</li> </ul>                     |
| DFaker                            | <a href="https://github.com/dfaker/df">https://github.com/dfaker/df</a>                                                           | <ul style="list-style-type: none"> <li>- DSSIM loss function [45] is used to reconstruct face.</li> <li>- Implemented based on Keras library.</li> </ul>                                                                                                |
| DeepFake_tf                       | <a href="https://github.com/StromWine/DeepFake_tf">https://github.com/StromWine/DeepFake_tf</a>                                   | Similar to DFaker but implemented based on tensorflow.                                                                                                                                                                                                  |
| AvatarMe                          | <a href="https://github.com/lattas/AvatarMe">https://github.com/lattas/AvatarMe</a>                                               | <ul style="list-style-type: none"> <li>- Reconstruct 3D faces from arbitrary “in-the-wild” images.</li> <li>- Can reconstruct authentic 4K by 6K-resolution 3D faces from a single low-resolution image [46].</li> </ul>                                |
| MarionETte                        | <a href="https://hyperconnect.github.io/MarionETte">https://hyperconnect.github.io/MarionETte</a>                                 | <ul style="list-style-type: none"> <li>- A few-shot face reenactment framework that preserves the target identity.</li> <li>- No additional fine-tuning phase is needed for identity adaptation [47].</li> </ul>                                        |
| DiscoFaceGAN                      | <a href="https://github.com/microsoft/DiscoFaceGAN">https://github.com/microsoft/DiscoFaceGAN</a>                                 | <ul style="list-style-type: none"> <li>- Generate face images of virtual people with independent latent variables of identity, expression, pose, and illumination.</li> <li>- Embed 3D priors into adversarial learning [48].</li> </ul>                |
| StyleRig                          | <a href="https://gvv.mpi-inf.mpg.de/projects/StyleRig">https://gvv.mpi-inf.mpg.de/projects/StyleRig</a>                           | <ul style="list-style-type: none"> <li>- Create portrait images of faces with a rig-like control over a pretrained and fixed StyleGAN via 3D morphable face models.</li> <li>- Self-supervised without manual annotations [49].</li> </ul>              |
| FaceShifter                       | <a href="https://lingzhili.com/FaceShifterPage">https://lingzhili.com/FaceShifterPage</a>                                         | <ul style="list-style-type: none"> <li>- Face swapping in high-fidelity by exploiting and integrating the target attributes.</li> <li>- Can be applied to any new face pairs without requiring subject specific training [50].</li> </ul>               |
| FSGAN                             | <a href="https://github.com/YuvalNirkin/fsgan">https://github.com/YuvalNirkin/fsgan</a>                                           | <ul style="list-style-type: none"> <li>- A face swapping and reenactment model that can be applied to pairs of faces without requiring training on those faces.</li> <li>- Adjust to both pose and expression variations [51].</li> </ul>               |
| Transformable Bottleneck Networks | <a href="https://github.com/kyleolsz/TB-Networks">https://github.com/kyleolsz/TB-Networks</a>                                     | <ul style="list-style-type: none"> <li>- A method for fine-grained 3D manipulation of image content.</li> <li>- Apply spatial transformations in CNN models using a transformable bottleneck framework [52].</li> </ul>                                 |
| “Do as I Do” Motion Transfer      | <a href="https://github.com/carolinec/EverybodyDanceNow">github.com/carolinec/EverybodyDanceNow</a>                               | <ul style="list-style-type: none"> <li>- Automatically transfer the motion from a source to a target person by learning a video-to-video translation.</li> <li>- Can create a motion-synchronized dancing video with multiple subjects [53].</li> </ul> |
| Neural Voice Puppetry             | <a href="https://justusthies.github.io/posts/neural-voice-puppetry">https://justusthies.github.io/posts/neural-voice-puppetry</a> | <ul style="list-style-type: none"> <li>- A method for audio-driven facial video synthesis.</li> <li>- Synthesize videos of a talking head from an audio sequence of another person using 3D face representation. [54].</li> </ul>                       |



**Figure 7**

From: [A prosocial fake news intervention with durable effects](#)



Bullshit receptivity immediately after the intervention (left panel) and 4 weeks later (right panel). Error bars represent standard errors. The y-axis represents standardized scores in terms of bullshit receptivity ratings controlled for pre-intervention bullshit receptivity ratings. Higher scores indicate beliefs in bullshit are meaningful such as "interdependence is rooted in ephemeral actions".

## A prosocial fake news intervention with durable effects

[Gábor Orosz](#) , [Benedek Paskuj](#), [Laura Faragó](#) & [Péter Krekó](#)

[Scientific Reports](#) **13**, Article number: 3958 (2023) | [Cite this article](#)

2544 Accesses | 2 Citations | 13 Altmetric | [Metrics](#)

### Abstract

The present online intervention promoted family-based prosocial values—in terms of helping family members—among young adults to build resistance against fake news. This preregistered randomized controlled trial study is among the first psychological fake news interventions in Eastern Europe, where the free press is weak and state-sponsored misinformation runs riot in mainstream media. In this intervention, participants were endowed with an expert role and requested to write a letter to their digitally less competent relatives explaining six strategies that help fake news recognition. Compared to the active control group there was an immediate effect ( $d = 0.32$ ) that persisted until the follow-up four weeks later ( $d = 0.22$ ) on fake news accuracy ratings of the young, advice-giving participants. The intervention also reduced the bullshit receptivity of participants both immediately after the intervention and in the long run. The present work demonstrates the power of using relevant social bonds for motivating behavior change among Eastern European participants. Our prosocial approach with its robust grounding in human psychology might complement prior interventions in the fight against misinformation.





This survey is part of a MIT scientific research project. Your decision to complete this survey is voluntary. Please do not provide any sensitive identifiable information that would allow your identity to readily be ascertained, directly or through identifiers provided in your responses. This study is investigating how humans and artificial intelligence algorithms interact. In the study, you will answer questions and have a back and forth discussion with an artificial intelligence algorithm. If you give us permission by saying yes below, we plan to discuss/publish the results in an academic forum. In any publication, information will be provided in such a way that you cannot be identified. Only members of the research team will have access to the original data set. Before the data is shared outside the research team, any potentially identifying information will be removed. Once identifying data has been removed, the data may be used by the research team, or shared with other researchers, for both related and unrelated research purposes in the future. Your anonymized data may also be made available in online data repositories such as the Open Science Framework, which allow other researchers and interested parties to use the data for further analysis. **Clicking on the button below indicates that you are at least 18 years of age and agree to complete this survey voluntarily.**

# DETEKTO TANANYAGOK

## ÁLHÍREK

Mik azok az álhírek és miért vagyunk rájuk fogékonyak...

1

Mik azok az “álhírek”?

2

Miért vagyunk  
fogékonyak  
az álhírekre?

3

Hogyan készülnek  
és terjednek  
a kamu képek  
és videók?

# PORNÓTÓL POLITIKÁIG: A DEEPPFAKE ITT VAN VELÜNK

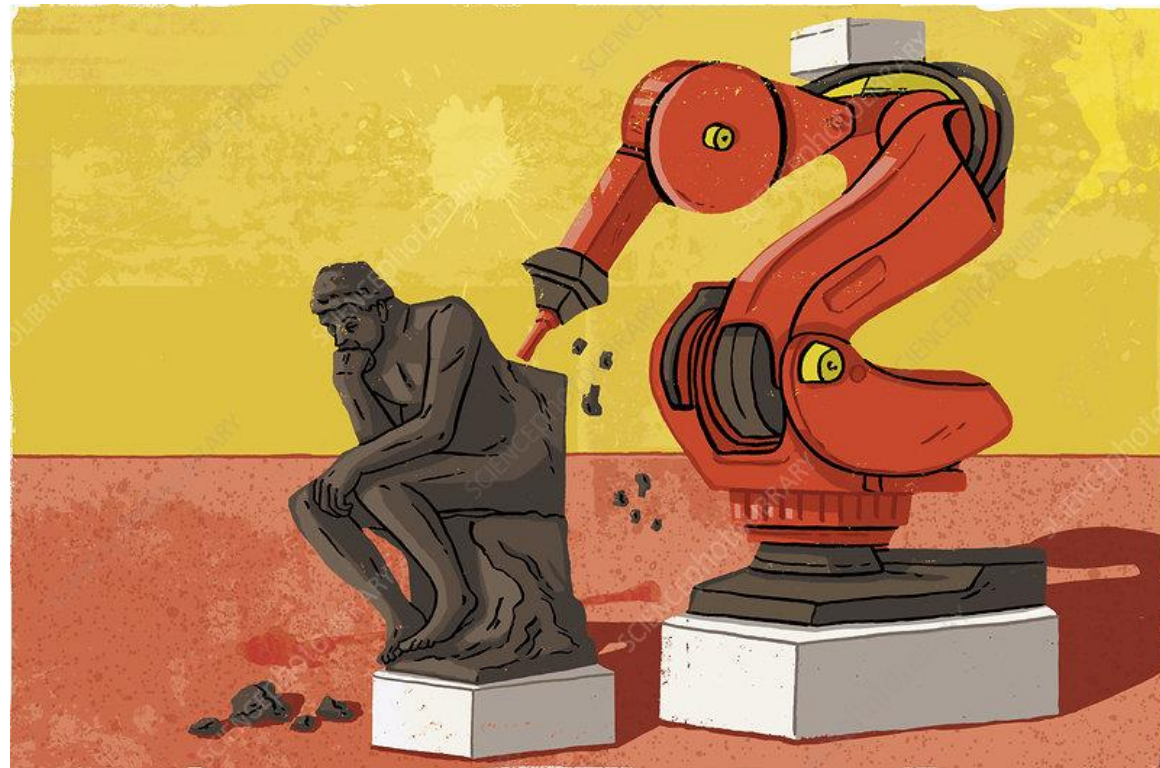
2022/06/30 - Lakmusz

*Diószegi-Horváth Nóra, Német Szilvi*

Mi a közös Scarlett Johanssonban, Leia hercegnőben és Vitalij Klicskóban? Mára ez nem egy közepesen rossz vicc kezdőkérdése, hanem a valóság: mindhármukól készült már élethű hamisított videó. Hol tart a deepfake, és hogyan szűrhetod ki a hamisítványt?

# Következtetések: főbb dilemmák, kihívások

- Demokratikus kihívás: pluralizmus helyett fragmentáció a politikai közösségben
- Információs tanult tehetetlenség (Nisbet&Kamenchuk, 2021); információs apokalipszis, valóság-apátia (Gregor & Mlejnková, 2021).
- Megerősítési torzítás és a morális érzelmek felkorbácsolása
- Emberkép – homo algorithmus. Szabad akarat, szabad választás?
- A kevesebb több? (digitális lábnyom, információfeldolgozás)
- Sem utópia, sem disztópia?







ELTE | PPK  
PEDAGÓGIAI ÉS PSZICHOLÓGIAI KAR



DEZINFORMÁCIÓ  
ÉS MESTERSÉGES  
INTELLIGENCIA  
KUTATÓCSOPORT

# Köszönöm a figyelmet!

Kreko.peter@ppk.elte.hu