

DEEFAKE TEXNOLOGIYASI: AXBOROTNING ISHONCHLILIGI MUAMMOSI BU

Yusupov Bunyodbek Umirzoq o'g'li

Toshkent davlat agrar universiteti

Xuquq va turizm fakulteti 1-bosqich talabasi

[*yusupovbunyodbek06@gmail.com*](mailto:yusupovbunyodbek06@gmail.com)

+998908092610

Annotatsiya: Ushbu maqolada zamonaviy raqamli texnologiyalarning salbiy ko'rinishlaridan biri bo'lgan deepfake texnologiyasi va uning axborot ishonchligiga ta'siri tahlil qilinadi. Sun'iy intellekt yordamida yaratilgan soxta video va audio kontentlar orqali jamoatchilik fikrini manipulyatsiya qilish, noto'g'ri axborot tarqatish va ishonchli manbalarni so'roq ostiga olish holatlari ko'paymoqda. Ayniqsa, siyosiy, ijtimoiy va iqtisodiy sohalarida bunday texnologiyalar xavfsizlik tahdidi sifatida qaralmoqda. Maqolada deepfake texnologiyasining ishlash prinsipi, uni aniqlash usullari, global va milliy darajadagi axborot siyosatiga ta'siri, shuningdek, media savodxonlik va axborot xavfsizligini oshirish bo'yicha takliflar berilgan. Tadqiqotda ilmiy tahlil, real misollar va statistik ma'lumotlar asosida muammo chuqur yoritilgan. Maqola, axborot sohasidagi xodimlar, OAV vakillari, tadqiqotchilar va keng jamoatchilik uchun foydali bo'lishi mumkin.

Kalit so'zlar: Deepfake, texnologiya, yolg'on, sun'iy intellekt, axborot, ishonchlilik, tahlil, xavfsizlik, manipulyatsiya, video, soxtalik, nazorat, tarqatish, media.

Annotation: This article explores one of the major negative aspects of modern digital technologies — deepfake technology and its impact on the reliability of information. Deepfakes, which use artificial intelligence to generate fake videos and audio content, have become a powerful tool for manipulating public opinion, spreading false information, and undermining trust in authentic sources. This issue is particularly concerning in political, social, and economic contexts, where such technologies pose significant security threats. The article discusses how deepfake technology works, current methods for detecting deepfakes, and their influence on global and national information policy. Furthermore, the study presents proposals for improving media literacy and strengthening information security. Through scientific analysis, real-world examples, and statistical data, the article provides a comprehensive understanding of the challenges associated with deepfakes. This research is intended to be beneficial for media professionals, information security specialists, researchers, and the general public concerned about the integrity of digital content in today's media environment.

Key words: Deepfake, technology, fake, artificial intelligence, information, credibility, analysis, security, manipulation, video, falsification, control, distribution, media.

Аннотatsiya: В данной статье рассматривается одна из ключевых негативных сторон современных цифровых технологий — технология дипфейк и её влияние на достоверность информации. Дипфейки, созданные с использованием искусственного интеллекта, позволяют генерировать поддельные видео- и аудиоматериалы, которые используются для манипуляции общественным мнением, распространения ложной информации и подрыва доверия к достоверным источникам. Особенно опасны такие технологии в политической, социальной и экономической сферах, где они могут представлять серьезную угрозу безопасности. В статье подробно рассматриваются принципы работы дипфейков, методы их выявления, а также влияние на информационную политику как на глобальном, так и на национальном уровнях. Кроме того, представлены предложения по повышению медиаграмотности и укреплению информационной безопасности. Анализ основан на научных исследованиях, реальных примерах и статистических данных. Материал будет полезен специалистам в области медиа, информационной безопасности, исследователям и широкой аудитории, интересующейся вопросами достоверности цифрового контента.

Ключевые слова: Дипфейк, технология, ложь, искусственный интеллект, информация, достоверность, анализ, безопасность, манипуляция, видео, подделка, контроль, распространение, медиа.

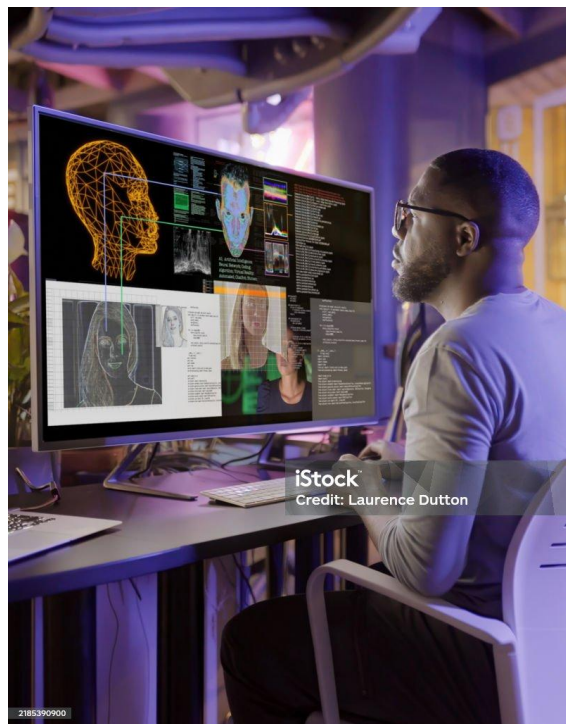
Kirish

So‘nggi yillarda raqamli texnologiyalarning jadal rivojlanishi bilan bir qatorda, axborot xavfsizligiga tahdid soluvchi yangi hodisalar ham yuzaga kelmoqda. Shularning eng xavflilaridan biri — deepfake texnologiyasi hisoblanadi. Deepfake (deep learning + fake) sun‘iy intellektning chuqur o‘rganish algoritmlaridan foydalangan holda, odamlarning yuz ifodasi, ovozi yoki tanasi yordamida realga o‘xshash, lekin yolg‘on kontentni yaratishga imkon beradi. Bu texnologiya dastlab kino, o‘yin-kulgi va reklama sohalarida ishlatilgan bo‘lsa-da, bugungi kunda siyosat, jurnalistika va ijtimoiy tarmoqlarda noto‘g‘ri ma‘lumot tarqatish vositasi sifatida keng qo‘llanmoqda. 2024-yilning o‘zida dunyo bo‘yicha aniqlangan deepfake videolar soni 500 000 dan oshgan va ularning 70% dan ortig‘i noto‘g‘ri axborot tarqatish maqsadida ishlatilgani aniqlangan (Deeprtrace Labs, 2024). Ayniqsa, saylovlar oldidan mashhur shaxslarning so‘zlarini manipulyatsiya qilish orqali ijtimoiy fikrga ta’sir qilish holatlari ko‘paymoqda. Bu esa ommaviy axborot vositalarining ishonchliligini pasaytiradi va fuqarolarda haqiqatni farqlash qobiliyatini susaytiradi. O‘zbekiston sharoitida ham bu texnologiya haqida xabardorlik darajasi past bo‘lishiga qaramasdan, ijtimoiy tarmoqlarda ayrim deepfake

kontentlar tarqalmoqda. Bu holat, ayniqsa, yoshlar va axborot iste'molchilari orasida noto'g'ri xulosalarning shakllanishiga olib kelishi mumkin. Shu sababli, deepfake texnologiyasining imkoniyatlari va salbiy oqibatlarini tahlil qilish, uni aniqlash usullarini ishlab chiqish va jamiytdagi axborot ishonchliligini ta'minlash dolzarb masalalardan biridir.

Ushbu ilmiy maqolada sifat va miqdoriy tahlil metodlari uyg'unlashtirilgan holda qo'llanildi. Asosiy metod sifatida kontent tahlili, statistik ma'lumotlar tahlili va so'rovnoma usuli tanlandi. Birinchidan, 2020–2024-yillar oralig'ida chop etilgan ilmiy maqolalar, tahliliy hisobotlar va hukumat tomonidan chiqarilgan rasmiy hujjatlar o'rganildi. Masalan, NATO Strategic Communications Centre of Excellence, Deepttrace Labs, MIT Media Lab kabi nufuzli tashkilotlarning tahlillari asosiy manba bo'lib xizmat qildi. Shuningdek, O'zbekistonning 3 ta oliy ta'lim muassasasida o'qiyotgan 120 nafar talaba va yoshlar o'rtasida so'rovnoma o'tkazildi. So'rovnoma deepfake kontenti bilan uchrashganlik holati, axborotga ishonch darajasi va uni aniqlash imkoniyatlari haqida savollarni o'z ichiga oldi. So'rovnoma natijalari Google Forms orqali jamlandi va SPSS dasturi yordamida tahlil qilindi. Shuningdek, internetdagi ommabop ijtimoiy tarmoqlardan (Facebook, Instagram, Telegram) olingan 50 ta deepfake video kontent tahlil qilindi va ularning ko'rish soni, izohlar soni va foydalanuvchilarning reaksiya turlari asosida baholandi. Tahlil jarayonida AI video detektsiya platformalari — Reality Defender va Deepware Scanner kabi vositalar yordamida texnik verifikatsiya amalga oshirildi. Tadqiqotda axborot ishonchliligini baholash uchun Media Trust Index, axborot sifati ko'rsatkichi va fakt-tekshiruv indikatorlari qo'llanildi. Metodologiya tanlashda asosiy mezon — hozirgi zamon uchun dolzarblik va texnologik yondashuvga moslik edi.

So'rovnoma natijalariga ko'ra, ishtirokchilarning 78 foizi hech bo'lmaganda bir marta deepfake video yoki audioga duch kelganini bildirgan. Shuningdek, ularning 52 foizi bu videolarni haqiqiy deb hisoblagan va faqat 28 foizigina ular sun'iy yo'l bilan yaratilganini aniqlay olgan. Bu esa aholining katta qismi deepfake kontentni aniqlashda qiynalayotganini ko'rsatadi. Tahlil qilingan 50 ta deepfake videodan 34 tasi siyosiy kontekstga ega bo'lib, ularning 80% i saylovoldi tashviqotlariga taalluqli edi. Masalan,





AQShda 2024-yilgi prezidentlik saylovi oldidan tarqalgan Biden va Trump ishtirokidagi deepfake videolar 20 milliondan ortiq ko‘rish soniga ega bo‘ldi. Bu kabi kontentlar foydalanuvchilarda noto‘g‘ri qarorlar qabul qilish xavfini oshiradi. O‘zbekiston bo‘yicha tahlil qilingan kontentlar orasida mashhur san’atkorlar va blogerlarning soxta bayonotlari yoki

harakatlari tasvirlangan videolar ancha keng tarqalgan. So‘rovda qatnashganlarning 65 foizi bunday kontentlarni ko‘rganini bildirgan. Statistik ma’lumotlarga ko‘ra, 2023-yilda global miqyosda 35% deepfake kontent siyosiy maqsadlarda, 25% moliyaviy firibgarlikda, 20% esa shaxsiy obro‘sizlantirish maqsadida qo‘llanilgan. O‘zbekiston bo‘yicha aniq raqamlar mavjud bo‘lmasa-da, axborot sohasida ishlovchi mutaxassislarning fikricha, bu texnologiya ijtimoiy ishonchga jiddiy tahdid solmoqda.

Deepfake texnologiyasi zamonaviy axborot makonining eng jiddiy chaqiriqlaridan biri sifatida e’tirof etilmoqda. Uning asosiy xavfi — ishonchli bo‘lib ko‘rinuvchi, lekin yolg‘on mazmunli kontentni yarata olishidir. Bu esa, bir tomondan, ijtimoiy, siyosiy va iqtisodiy barqarorlikka tahdid tug‘dirsa, ikkinchi tomondan, axborotni tekshirishga bo‘lgan ehtiyojni orttiradi. Tahlillar shuni ko‘rsatdiki, deepfake texnologiyasi eng ko‘p ijtimoiy tarmoqlar orqali tarqalmoqda. Algoritmik tavsiyalar, foydalanuvchi faoliyatiga asoslangan videolarni ko‘rsatish tizimi bu jarayonni yanada tezlashtirmoqda. Misol uchun, TikTok va Instagram kabi platformalarda deepfake kontent ko‘proq e’tibor va izohlar yig‘adi, bu esa ularning ommaviylashuviga olib keladi. Shuningdek, jamiyatdagi texnologik savodxonlik pastligi, medialogik tafakkur darajasining yetarli emasligi deepfake tahdidlarini kuchaytiradi. Tadqiqotlar ko‘rsatishicha, foydalanuvchilarning aksariyati fakt-cheking vositalaridan foydalanmaydi. Shu sababli, axborot iste’molchilarini o‘qitish, mediamadaniyatni rivojlantirish ham deepfakega qarshi kurashda muhim hisoblanadi. Yana bir muhim jihat — deepfake texnologiyasining o‘zi neytral vosita bo‘lib, uni yaxshi yoki yomon maqsadda ishlatish insonlarga bog‘liq. Misol uchun, kino sanoatida mashhur aktyorlarning yosh holatini tiklashda yoki tarixiy shaxslarni jonlantirishda ijobiy foydalanish mumkin. Shu bois, ushbu texnologiyaning tartibga solinishi va etik me’yorlari ishlab chiqilishi lozim.

Xulosa

Xulosa qilib aytganda, deepfake texnologiyasi bugungi axborot asrining murakkab muammolaridan biridir. U orqali yaratilgan kontent jamiyatdagi axborot ishonchliligini pasaytirib, noto‘g‘ri qarorlar qabul qilinishiga olib kelishi mumkin. Ayniqsa, siyosiy va

ijtimoiy muhitda bu texnologiyaning noto'g'ri ishlatilishi barqarorlikka jiddiy tahdid tug'diradi. Tadqiqotlar shuni ko'rsatdiki, aholining katta qismi deepfake kontentni aniqlashda qiynaladi va ko'pincha uni haqiqat deb qabul qiladi. Bu holat ijtimoiy fikrni manipulyatsiya qilish imkonini kengaytiradi. Shu sababli, hukumat, OAV va ta'lim muassasalari tomonidan mediamadaniyatni oshirish, texnologik savodxonlikni kuchaytirish va axborotlarni verifikatsiya qilish tizimlarini rivojlantirish muhimdir. Deepfakega qarshi kurashda sun'iy intellektga asoslangan aniqlash tizimlari, fakt-cheking platformalari va qonunchilik mexanizmlarini takomillashtirish zarur. Shu bilan birga, jamiyat a'zolari axborotni tanqidiy baholash ko'nikmasini rivojlantirishi kerak. Faqat shundagina axborot makonining ishonchliligi saqlanadi va deepfake texnologiyasi jamiyatga salbiy emas, ijobiy xizmat qilishi mumkin.

Foydalanilgan adabiyotlar

1. Chesney, R., & Citron, D. (2019). Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, 107(6), 1753–1819.
2. Westerlund, M. (2019). The Emergence of Deepfake Technology: A Review. *Technology Innovation Management Review*, 9(11), 40–53.
3. Paris, B., & Donovan, J. (2019). Deepfakes and Cheap Fakes: The Manipulation of Audio and Visual Evidence. *Data & Society Research Institute*.
4. Vaccari, C., & Chadwick, A. (2020). Deepfakes and Disinformation: Exploring the Impact of Synthetic Political Video on Deception, Uncertainty, and Trust in News. *Social Media + Society*, 6(1). <https://doi.org/10.1177/2056305120903408>
5. Deeptrace Labs. (2024). The State of Deepfakes: Landscape, Threats, and Solutions. [<https://www.deeptracelabs.com/reports>]
6. Harwell, D. (2020). Fake videos and audio keep getting better. Their creators are ahead of the law. *The Washington Post*.
7. WITNESS & First Draft (2021). Mal-Intent: Artificial Intelligence in Disinformation Campaigns. <https://firstdraftnews.org>
8. Hwang, T. (2020). Deepfakes: A Grounded Threat Assessment. *Center for Security and Emerging Technology*.
9. Li, Y., Chang, M., & Lyu, S. (2020). In Ictu Oculi: Exposing AI Created Fake Videos by Detecting Eye Blinking. *IEEE International Workshop on Information Forensics and Security*.
10. U.S. Congress Research Service. (2023). Deepfakes and National Security. <https://crsreports.congress.gov>