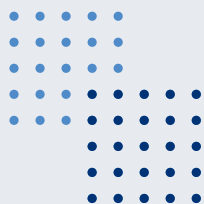


#Halo! Tu cyberbezpieczny Senior!

Nie zadzierAI ze mną



NASK



Ministerstwo
Cyfryzacji



PROJEKT FINANSOWANY ZE ŚRODKÓW
MINISTERSTWA CYFRYZACJI

#Halo! Tu cyberbezpieczny Senior!

Nie zadzierAI ze mną

Beata Frankiewicz

Sylwia Adamczyk

#Halo! Tu cyberbezpieczny Senior!
Nie zadzieraj ze mną

Autorki: **Beata Frankiewicz, Sylwia Adamczyk**

Konsultacja merytoryczna: **Karol Bojke**

Redakcja językowa: **Łukasz Szczęsny, Agnieszka Filipkowska**

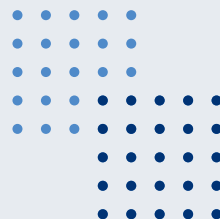
Opracowanie graficzne: **NASK – PIB**

Publikacja jest rozpowszechniana na zasadach licencji Creative Commons Uznanie autorstwa – Użycie niekomercyjne (CC BY-NC) 4.0 Międzynarodowe

ISBN: 978-83-68356-42-7

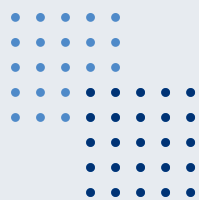
NASK – Państwowy Instytut Badawczy
ul. Kolska 12, 01-045 Warszawa
www.nask.pl

2026



Publikacja wyraża jedynie poglądy autorów i nie może być utożsamiana z oficjalnym stanowiskiem Ministra Cyfryzacji.

Spis treści



Czym jest sztuczna inteligencja?	5
„AI generated” – co to oznacza?	6
Ukryta siła technologii – gdzie spotykasz AI?	7
Twój smartfon i AI – pomoc, której nie widać	7
AI w obsłudze klienta – pomoc bez czekania	8
AI w wyszukiwarkach internetowych – odpowiedzi w mgnieniu oka	8
Chaty AI – szansa czy pułapka?	9
AI – pomocne, ale czasem niebezpieczne. Dlaczego trzeba uważać?	10
Media społecznościowe i „papka AI”	12
Wzruszający obrazek? Uważaj, to może być podstęp!	16
Deepfake w dezinformacji i cyberprzestępstwach	18
Próba oszustwa? Zgłoś ją	21



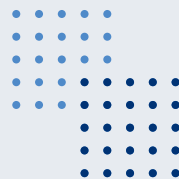
Czym jest sztuczna inteligencja?



Sztuczna inteligencja to technologia komputerowa, która wykorzystuje ogromne ilości danych do wykonywania pewnych zadań, które kiedyś potrafił wykonywać tylko człowiek. Narzędzia oparte na sztucznej inteligencji mogą m.in. generować teksty, realistyczne obrazy czy wideo lub odpowiadać na pytania zadane w codziennym języku.

Rozwój sztucznej inteligencji może mieć pozytywny wpływ na różne dziedziny życia. Technologia ta jest wykorzystywana w biznesie czy medycynie i ma wiele pożytecznych zastosowań. Jednak wraz z korzyściami płynącymi z jej rozwoju pojawiają się pewne wyzwania i zagrożenia.

Nie każdy używa sztucznej inteligencji w dobrych celach. Oparte na niej narzędzia mogą być wykorzystywane do oszustw i manipulacji. Na kolejnych stronach dowiesz się, jak je rozpoznawać i jak się przed nimi chronić.



„AI generated” – co to oznacza?

Skrót „AI” pochodzi z języka angielskiego, od słów „**artificial intelligence**”, tłumaczonych dokładnie jako „**sztuczna inteligencja**”. Coraz częściej można spotkać ten skrót w internecie – na stronach czy w aplikacjach.

Czasami w sieci można natknąć się także na oznaczenie „AI generated”, czyli „wygenerowane przez sztuczną inteligencję” albo „AI generated image” – „obraz wygenerowany przez sztuczną inteligencję”. Treści z takim napisem – na przykład zdjęcia, filmy lub teksty – zostały stworzone przez komputer, a nie przez człowieka. Takie oznaczenie może wyglądać tak jak na obrazku poniżej.

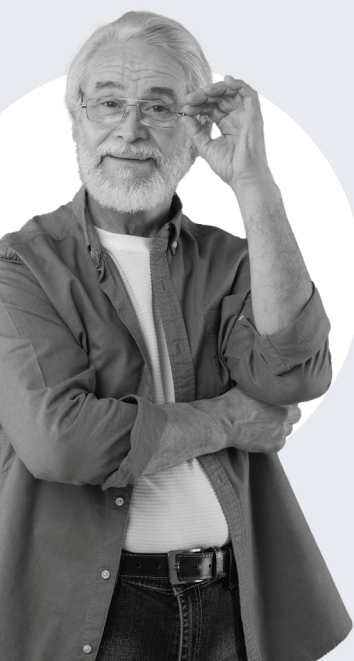


„AI-Generated Image”



ŹRÓDŁO: FACEBOOK.COM

Niestety nie każda sztucznie wygenerowana treść jest oznaczana w ten sposób. Dlatego też warto zachować czujność. Trzeba pamiętać, że nawet realistycznie wyglądające zdjęcie czy nagranie nie zawsze przedstawia prawdziwe wydarzenia. Więcej na ten temat przeczytasz w rozdziale o deepfejkach (**STRONA 18**).



Ukryta siła technologii – gdzie spotykasz AI?



Kiedy myślimy o sztucznej inteligencji, wielu z nas wyobraża sobie skomplikowane urządzenia, laboratoria i technologie, które są „gdzieś daleko”. Tymczasem AI stało się częścią naszej codzienności szybciej, niż się spodziewaliśmy. Wiele działań, które wykonujemy każdego dnia, jest obsługiwanych w tle przez systemy uczące się na podstawie ogromnych zbiorów danych.

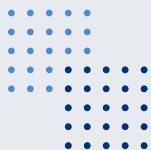
Najłatwiej dostrzec to w przypadku telefonów, które mamy zawsze pod ręką. Korzystamy z nich naturalnie, nie zastanawiając się nad tym, że w tle pracują mechanizmy, które ułatwiają nam życie – podpowiadają, poprawiają i pomagają, zanim zdążymy o to poprosić.

Twój smartfon i AI – pomoc, której nie widać

Smartfon to dziś coś więcej niż telefon – to mały komputer, który uczy się naszych nawyków i pomaga w codziennych zadaniach.

- **Pisanie wiadomości** – telefon nie tylko poprawia literówki i podpowiada kolejne wyrazy. Z czasem uczy się naszego stylu pisania, zapamiętuje najczęściej używane zwroty i dopasowuje sugestie tak, aby brzmiały naturalnie, jak nasze własne słowa.
- **Fotografia** – aparat sam dobiera parametry, m.in. jasność, ostrość i kolory, by zdjęcie wyglądało lepiej.
- **Polecenia głosowe** – po uruchomieniu odpowiedniej funkcji lub aplikacji wystarczy wypowiedzieć polecenie, a system zamieni mowę na tekst lub wykona zadanie.





To, co kiedyś wymagało wiedzy i umiejętności, dziś dzieje się automatycznie – w tle, bez naszej uwagi. To ogromne ułatwienie, które oszczędza czas i wysiłek. Jednak im bardziej polegamy na automatyzacji, tym trudniej zauważyć, gdzie kończy się nasz wybór, a zaczyna decyzja algorytmu. Warto o tym pamiętać, bo wygoda nie zawsze idzie w parze z pełną kontrolą.

AI w obsłudze klienta – pomoc bez czekania



Wiele osób może nie zdawać sobie sprawy, że często na infoliniach najpierw rozmawiamy z wirtualnym doradcą, a dopiero później z prawdziwym konsultantem. Firmy sięgają po takie rozwiązania, bo wiele pytań powtarza się każdego dnia: „Gdzie jest moja paczka?”, „Jak zmienić hasło?”, „Dlaczego moja faktura jest wyższa?”.

Abyśmy nie musieli czekać w kolejce na infolinię i słuchać muzyki, sprawę przejmuje wirtualny doradca. Analizuje pytanie i korzysta z ogromnej bazy wcześniejszych rozmów, żeby podać gotową odpowiedź. Jeśli problem jest bardziej skomplikowany, rozmowa trafia do prawdziwego konsultanta. Dzięki temu, że na proste pytania odpowiada automat, pracownicy mają czas na trudniejsze sprawy.

Firmy chętnie korzystają z pomocy AI, bo działa błyskawicznie, jest dostępne 24 godziny na dobę przez siedem dni w tygodniu i obsługuje setki osób jednocześnie. Efekt? Większość spraw możemy załatwić od ręki, a te wymagające większej uwagi szybciej trafiają do specjalisty. Firmy oszczędzają czas, a klienci – nerwy związane z czekaniem na obsługę.

AI w wyszukiwarkach internetowych – odpowiedzi w mgnieniu oka



Gdy za pomocą wyszukiwarki internetowej zaczynamy formułować pytanie, AI podpowiada kolejne słowa, bazując na historii naszych dotychczasowych zapytań.

Oprócz tego często pierwsza odpowiedź wyświetlająca się na górze strony to wynik działania algorytmu, który automatycznie przygotowuje

Zastanawiasz się, jak to działa? Wpisz w wyszukiwarkę pytanie, np. „jak pozbyć się plamy z kawy?”. AI przeanalizuje strony opisujące dany temat, wybierze najważniejsze fakty i ułoży je w prosty tekst zamieszczony na początku listy wyników. Zajmie to tylko kilka sekund. Oto przykład:



Aby usunąć plamy z kawy, działaj szybko, używając domowych sposobów jak pasta z sody oczyszczonej i wody, roztwór octu z wodą, sok z cytryny lub sól, delikatnie pocierając, a następnie wypłucz i upierz tkaninę, pamiętając o zasadzie: nie używaj gorącej wody na świeże plamy i unikaj mocnego tarcia, które utrwala zabrudzenie.



- **Natychmiast:** Polej obficie gorącą wodą (jeśli to dozwolone dla materiału) lub nasacz chusteczkami nawilżającymi.

Pokaż więcej 

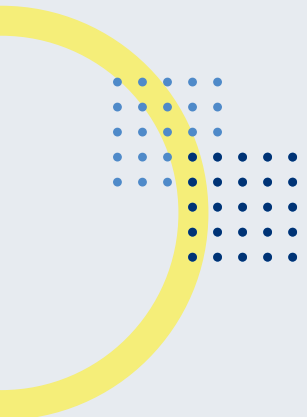
ŹRÓDŁO: GOOGLE.COM

Podsumowanie jest często trafne i oszczędza czas, zwłaszcza gdy odpowiada na proste pytania. Czasem jednak zawiera błędy. Dlatego warto upewniać się co do faktów, sięgając do innych źródeł – zwłaszcza w przypadku istotnych zagadnień, np. tych dotyczących zdrowia.

Chaty AI – szansa czy pułapka?

AI nie jest nieomylny – tworzy odpowiedzi na podstawie dostępnych danych, ale nie sprawdza ich poprawności. W efekcie zdarza się, że zwraca informacje, które wyglądają wiarygodnie, choć są fałszywe. Takie błędy nazywa się halucynacjami AI. To sytuacje, gdy sztuczna inteligencja podaje nieprawdziwe informacje, wymyśla porady zdrowotne czy prawne, powołuje się na nieistniejące źródła i badania naukowe. Co gorsza, robi to w przekonujący sposób, używając specjalistycznych słów, jakby wszystko było prawdą.





Korzystanie z narzędzi AI, takich jak czaty, może być naprawdę pomocne. Dzięki nim łatwiej zrozumieć trudne pojęcia, napisać pismo urzędowe, znaleźć kluczowe informacje w artykule, znaleźć przepis na obiad z produktów, które mamy akurat w lodówce, czy dowiedzieć się, jak załatwić sprawę online. Ich popularność rośnie, bo realnie zmieniają sposób, w jaki pracujemy i żyjemy – upraszczają codzienne zajęcia, przyspieszają pracę i pomagają szybko radzić sobie z problemami, których rozwiązanie kiedyś mogło zajmować wiele czasu. Jednak nawet najbardziej zaawansowane algorytmy nie zastąpią naszego krytycznego myślenia. To my musimy ocenić, czy odpowiedź jest prawdziwa, bezpieczna i zgodna z kontekstem – AI nie może wziąć odpowiedzialności za decyzje, które będziemy podejmować na bazie dostarczonych przez nie informacji.

AI – pomocne, ale czasem niebezpieczne. Dlaczego trzeba uważać?

Zdarza się, że czaty AI podają informacje nieprawdziwe, a nawet groźne dla zdrowia. Polskie media opisywały między innymi następujące przykłady:

- **Badania włoskich naukowców** pokazały, że w kwestiach medycznych AI myli się w aż 70% przypadków, a 30% odwołań do źródeł było niedokładnych lub całkowicie zmyślonych¹.
- **Pewien 60-latek trafił do szpitala** po tym, jak chciał wyeliminować sól z diety i zastąpił ją bromem – zgodnie z poradą przygotowaną przez chat AI².

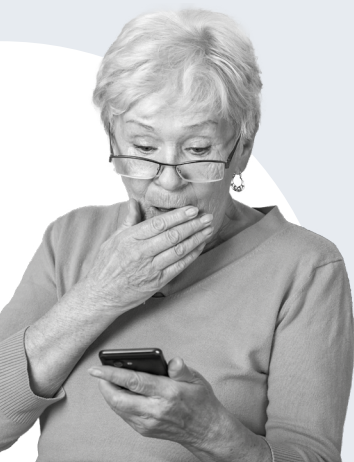


Sytuacja, gdy AI powołuje się na badania naukowe, które nie istnieją, nie jest niczym rzadkim. Sztuczna inteligencja potrafi wymyślić nazwiska lekarzy, tytuły publikacji i instytuty badawcze. Polska dziennikarka specjalizująca się w temacie zdrowia i medycyny opisywała przypadek, w którym AI cytowało „badania”, brzmiące bardzo wiarygodnie, choć były całkowicie fikcyjne, a potem z równie dużym zaangażowaniem przyznawało jej rację, gdy zarzuciła mu oszustwo³.

1 <https://naukawpolsce.pl/aktualnosci/news%2C110089%2Cwlochy-badania-sztuczna-inteligencja-uzywana-w-medycynie-popelnia-bledy-w>

2 <https://www.medonet.pl/zdrowie/wiadomosci,zatrul-sie-dieta-ulozona-przez-internet--tak-ai-niemal-doprowadzila-do-smierci-60-latka,artykul,49467426.html>

3 <https://swiatlekarza.pl/jak-oszukala-mnie-sztuczna-inteligencja/>



Dlaczego to groźne? Bo daje fałszywe poczucie bezpieczeństwa. Niektórzy rezygnują z wizyty u lekarza czy dietetyka, ignorują objawy i polegają na poradach z czatu. To może opóźnić leczenie albo doprowadzić do poważnych powikłań.

AI może pomóc w wyjaśnianiu trudnych pojęć, streszczaniu artykułów czy udzielaniu porad technicznych. Nie jest jednak narzędziem do diagnozowania, leczenia ani oceny stanu zdrowia. Nie zastąpi lekarza, dietetyka ani farmaceuty.



W sprawach zdrowia:

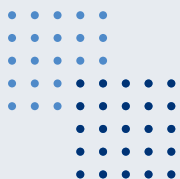
- Jeśli coś Cię niepokoi – skontaktuj się z lekarzem.
- Jeśli chcesz zmienić dietę – porozmawiaj ze specjalistą.
- Jeśli AI zasugeruje suplement diety lub lek – nie stosuj go na własną rękę.
- Jeśli AI powołuje się na badania – sprawdź informacje o nich w kilku źródłach i upewnij się, czy naprawdę ktoś je przeprowadził.



Zanim zaufasz poradzie z czatu, zastanów się, czy dotyczy ona spraw, w których potrzebna jest fachowa wiedza – na przykład zdrowia, prawa czy finansów. W takich sytuacjach AI może być pomocne jako narzędzie, ale nigdy nie powinno zastępować specjalisty.



Media społecznościowe i „papka AI”

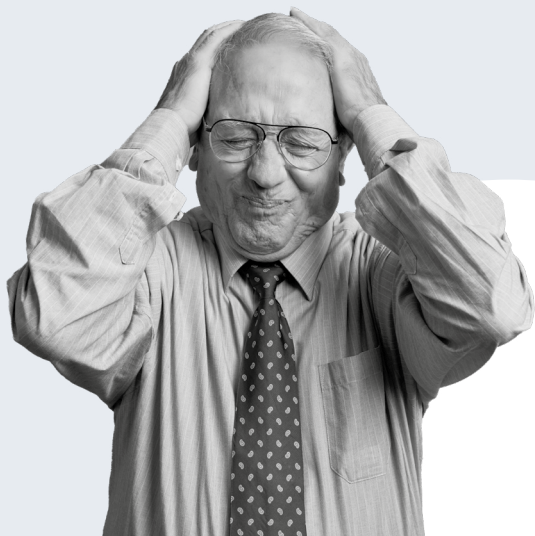


Na różne „wytwory” sztucznej inteligencji możemy natknąć się w mediach społecznościowych, takich jak Facebook, Instagram czy TikTok. Część udostępnianych tam materiałów została wygenerowana dzięki narzędziom AI.

Jedną z pułapek, z którą możemy mieć tutaj do czynienia, jest tzw. „AI slop”. Słowo „slop” w języku angielskim oznacza dosłownie „pomyje”, natomiast termin ten tłumaczony jest na polski także jako „papka AI”. „Pomyje” sztucznej inteligencji to niskiej jakości treści, najczęściej występujące w formie obrazków przypominających prawdziwe zdjęcia lub filmów, które na pierwszy rzut oka wyglądają realistycznie. Mają one za zadanie zachęcić nas do reakcji: polubienia danego wpisu, napisania komentarza lub udostępnienia.



Takie treści zazwyczaj bazują na naszych emocjach, przedstawiając sytuacje wzbudzające nasze współczucie, nostalgię, powodujące smutek czy wzruszenie.



Do popularnych wątków należą:

Życie na wsi



„Jeśli nie gardzisz nami, że jesteśmy rolnikami, możesz się z nami przywitać”

Jeśli nie gardzisz nami, że jesteśmy rolnikami, możesz się z nami przywitać



ŹRÓDŁO: FACEBOOK.COM

Ta rodzina rolników w rzeczywistości nie istnieje. To nie zdjęcie – to obrazek wygenerowany dzięki narzędziom sztucznej inteligencji.



Niedocenione dzieła

„Mój tata i ja robiliśmy to razem, ale nikt tego nie doceniał!”

Mój tata i ja robiliśmy to razem, ale nikt tego nie doceniał!



ŹRÓDŁO: FACEBOOK.COM

Nie istnieje także ta imponująca rzeźba śnieżna. Tak naprawdę nigdy nie została wybudowana. Ten obrazek również nie jest prawdziwym zdjęciem.



Samotne urodziny

„Dziś mam 85 lat. Nie mam rodziny, dzieci, mieszkam sama. Zrobiłam sobie ciasto, ale nikt mi nawet pogratulował.”

Dziś mam 85 lat. Nie mam rodziny, dzieci, mieszkam sama. Zrobiłam sobie ciasto, ale nikt mi nawet nie pogratulował. 🥺💔



ŹRÓDŁO: FACEBOOK.COM

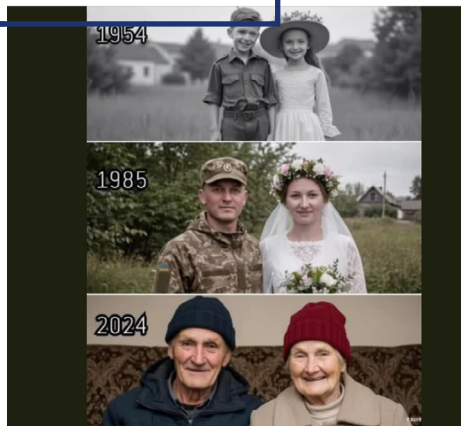
Również historia tej kobiety została zmyślona, a jej wizerunek szucznie wygenerowany.



Rocznice

„70 lat razem, co za wspaniała i szczęśliwa para”

70 lat razem, co za wspaniała i szczęśliwa para ❤️❤️



ŹRÓDŁO: FACEBOOK.COM

Wzruszające, ale ponownie – nieprawdziwe. W rzeczywistości żadna z tych osób nie istnieje.

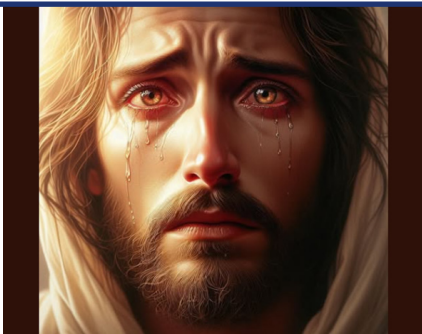


Religia



„Jezus płacze nad całą ludzkością i za każdym razem, gdy odrzucamy Jego miłość. Jeśli wierzysz, że ci, którzy kochają Boga, są najbardziej błogosławionymi ludźmi na świecie, kliknij link wyżej, a Boże błogosławieństwa spłyną na ciebie natychmiast.”

Jezus płacze nad całą ludzkością i za każdym razem, gdy odrzucamy Jego miłość. ❤️ Jeśli wierzysz, że ci, którzy kochają Boga, są najbardziej błogosławionymi ludźmi na świecie, kliknij link wyżej. 🙏, a Boże błogosławieństwa spłyną na ciebie natychmiast! 🙏❤️



ŹRÓDŁO: FACEBOOK.COM

W tym przypadku wizerunek płaczącego Jezusa ma poruszać osoby wierzące i skłonić je do interakcji. **Ten obrazek również wygenerowano dzięki AI.**



Nawiązania do bieżących wydarzeń



To zdjęcie przejdzie do historii!



Szacunek!

ŹRÓDŁO: FACEBOOK.COM

Za pomocą tego obrazka w czasie powodzi wyłudzano reakcje wdzięcznych pracującym bez wytchnienia służbom internautów. **Przedstawieni tu strażacy w rzeczywistości nie istnieją.**

Wzruszający obrazek? Uważaj, to może być podstęp!

Dzięki wykorzystaniu sztucznej inteligencji prowokacyjne grafiki można tworzyć łatwo i szybko. Nie wymaga to dużych zasobów ani żadnych umiejętności artystycznych.

Obrazek z silnym ładunkiem emocjonalnym ma na początku przyciągnąć uwagę użytkowników, którą oszuści mogą wykorzystać w nieuczciwe sposoby takie jak⁴:

- **Kradzież konta w mediach społecznościowych** – we wpisach towarzyszących obrazkom mogą się pojawiać linki prowadzące do fałszywych stron logowania, gdzie użytkownik nieświadomie podaje dane dostępu do swojego konta w mediach społecznościowych. Zamiast do zaufanego serwisu trafiają one w ręce przestępców. Taki profil jest przez nich przejmowany i wykorzystywany do dalszych oszustw. Sprawcy mogą na przykład wysłać z naszego konta wiadomości do naszych znajomych z prośbą o przesłanie pieniędzy.



Aby uchronić się przed przejęciem konta w internecie, warto włączyć dwuskładnikowe uwierzytelnianie (tzw. 2FA). Jest to dodatkowe zabezpieczenie, które polega na tym, że po wpisaniu hasła trzeba jeszcze potwierdzić logowanie w inny sposób – na przykład przepisując kod otrzymany SMS-em lub zatwierdzając powiadomienie w aplikacji na telefonie. Dzięki temu nawet jeśli ktoś pozna nasze hasło, nie będzie mógł zalogować się do danej platformy bez dostępu do naszego telefonu.



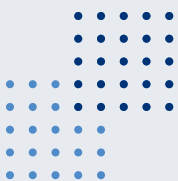
Włączenie dwuskładnikowego uwierzytelniania jest bardzo proste. W ustawieniach bezpieczeństwa konta (np. poczty elektronicznej, mediów społecznościowych itd.) wystarczy znaleźć opcję „Weryfikacja dwuetapowa”, „Uwierzytelnianie dwuskładnikowe” lub „Bezpieczeństwo logowania” i postępować zgodnie z wyświetlanymi instrukcjami. To niewielki wysiłek, który znacząco zwiększa bezpieczeństwo naszych danych i chroni nas przed oszustami.

⁴ <https://www.nask.pl/aktualnosci/zbyt-piekne-by-bylo-prawdziwe-emocjonalne-oszustwa-narzedziem-phishingu>

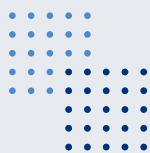
- **Wyłudzenie środków** – wpis z emocjonalnym obrazkiem, który zdobył dużą popularność, może zostać zmieniony – oszuści mogą zamieścić w nim zupełnie nową treść, na przykład link do rzekomej zbiórki charytatywnej. Wówczas całość przesłanych na taką „zbiórkę” pieniędzy trafia na konto przestępców.
- **Zbieranie danych do dalszych działań oszustów** – przez zamieszczane przy takich wpisach linki prowadzące do formularzy, fałszywych konkursów czy testów oszuści mogą pozyskiwać e-maile, numery telefonów, dane o lokalizacji i demograficzne. Mogą także kontaktować się w wiadomościach prywatnych z osobami, które zostawiają komentarze przy takich wpisach. W ten sposób rozbudowują swoją bazę potencjalnych kolejnych ofiar.
- **Dezinformacja** – po zdobyciu zaufania użytkowników – ich komentarzy i polubień, które przekładają się na popularność i wiarygodność konta – dany profil może zacząć promować treści polityczne, propagandowe czy pseudonaukowe. Dzięki zdobytej wcześniej rozpoznawalności, ten fałszywy przekaz będzie od razu trafiał do dużego grona obserwujących.



Dezinformacja to fałszywa bądź wprowadzająca w błąd treść, która jest tworzona lub rozpowszechniana celowo, z zamiarem oszukania odbiorców bądź uzyskania określonych korzyści, np. ekonomicznych lub politycznych. Podmioty, które rozpowszechniają dezinformację, często wykorzystują do tego sztuczną inteligencję – np. technologię deepfake, o której przeczytasz poniżej.



Deepfake w dezinformacji i cyberprzestępcach



Deepfake (ang. „głębokie fałszerstwo” – od „głębokiego uczenia”, metody rozwijania AI) to rodzaj technologii, która pozwala na tworzenie bardzo realistycznych, fałszywych zdjęć, filmów lub nagrań głosowych przy użyciu sztucznej inteligencji.

Generując deepfake, można sprawić, że czyjaś twarz lub głos zostaną ukazane w sytuacji, która nie wydarzyła się naprawdę.

Deepfake w dezinformacji

Deepfake może być narzędziem do szerzenia fałszywych informacji i manipulowania opinią publiczną. Wygenerowane za pomocą tej technologii zdjęcie czy nagranie może pokazywać np. osoby publiczne w kontekście, w którym w rzeczywistości nigdy się nie znalazły. Tak właśnie stało się w przypadku poniższego deepfajka:



W mediach społecznościowych pojawiło się „zdjęcie” wykonane rzekomo w dniu spotkania Donalda Trumpa z europejskimi przywódcami. Na wygenerowanym obrazku przedstawiono ich siedzących przed Gabinetem Ovalnym ze spuszczoneymi głowami, sugerując ich uległość. Fałsz można było jednak rozpoznać po szczegółach – niektóre postacie miały nienaturalną liczbę nóg.

Dzięki technologii deepfake można też generować nieistniejące w rzeczywistości postaci, tak jak widzimy na poniższym przykładzie (jest to kadr z dłuższego materiału wideo):



ŹRÓDŁO: X.COM

W mediach społecznościowych pojawiła się fałszywa informacja o tym, jakoby żona Wołodymyra Zełenskiego zakupiła luksusowe auto za prawie 4,5 mln euro. Dowodem na to miało być wideo nagrane przez pracownika paryskiego salonu samochodowego. Był to deepfake. Pracownik przedstawiony w wygenerowanym materiale w rzeczywistości nie istnieje. Producent zdementował informację o zakupie samochodu przez Olenę Zeleńską.

W mediach społecznościowych deepfejkі mogą rozprzestrzeniać się bardzo szybko, bo wyglądają autentycznie, a ludzie udostępniają je bez sprawdzania. Warto weryfikować informacje w różnych źródłach i nie podawać dalej niesprawdzonych treści.



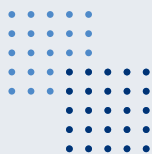
Deepfake w cyberprzestępcstwach

Technologia deepfake wykorzystywana jest też do oszustw internetowych.

Pewna Francuzka przekazała ok. 800 tys. euro oszustowi udającemu znanego aktora Brada Pitta. Była przekonana, że prowadzi romans z gwiazdą. Tymczasem oszust wyłudzał od niej środki finansowe, m.in. na rzekome opłaty celne od luksusowych prezentów, które miał jej wysłać. Przesyłał jej przy tym fałszywe zdjęcia i filmy z wyznaniem miłości, wygenerowane prawdopodobnie przy użyciu narzędzi AI, i przekonywał, że ma problemy zdrowotne:



ŹRÓDŁO: X.COM



Deepfejkki bywają również wykorzystywane w oszustwach inwestycyjnych. Oszuści tworzą filmy reklamowe, bezprawnie wykorzystując wizerunki znanych osób, np. prezydenta, ministra finansów, sportowców, dziennikarzy i biznesmenów. Ich sztucznie wygenerowane postaci „polecają” fałszywe inwestycje, obiecując pewny i wysoki zysk. Reklamy pojawiają się np. na Facebooku czy YouTube. Często wyglądają jak wywiady w telewizji. Gwiazda mówi: „Inwestuj teraz, to bezpieczne!”. Użytkownik klika link, zakłada konto na fałszywej platformie, wpłaca pieniądze, a potem nie może ich już wypłacić.



Pamiętaj: nie ufaj reklamom w mediach społecznościowych z celebrytami obiecującymi szybkie zyski. To może być oszustwo.

Próba oszustwa? Zgłoś ją



Widzisz podejrzaną reklamę, deepfejk lub fałszywą informację?
Zgłoś ją do NASK.

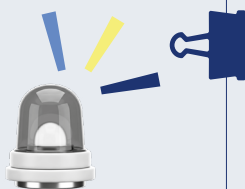


Incydenty takie jak np. podejrzanе wiadomości e-mail, deepfejki służące wyłudzeniu danych czy pieniędzy zgłaszaj do **CERT Polska** przez stronę: incydent.cert.pl lub przez aplikację **mObywatel** w zakładce „Bezpiecznie w sieci”.

SMS-y z podejrzanyimi linkami możesz zgłaszać, przekazując je na **bezpłatny numer 8080** za pomocą opcji „Przełącz dalej” w telefonie.

Treści o charakterze dezinformacyjnym i materiały zawierające fałszywe lub zmanipulowane informacje zgłaszaj do **Ośrodka Analizy Dezinformacji NASK** przez stronę: zglos-dezinformacje.nask.pl lub w wiadomości wysłanej bezpośrednio na adres mailowy: dezinformacja@nask.pl.

Dodatkowo, jeśli podejrzewasz, że dany materiał jest deepfejkiem, możesz zgłosić go do weryfikacji na adres: deepfake@nask.pl.



Oszustwa zgłaszaj bezpośrednio na policję.

Utratę danych dostępowych do bankowości lub kradzież środków z konta **zgłoś także do swojego banku.**

Pamiętaj!



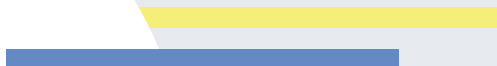
Gdy spotkasz kogoś, kto padł ofiarą oszustwa internetowego, lub przeczytasz o takiej historii w internecie, nie oceniaj tej osoby ani nie wytykaj jej naiwności – zarówno na żywo jak i w komentarzach w sieci.

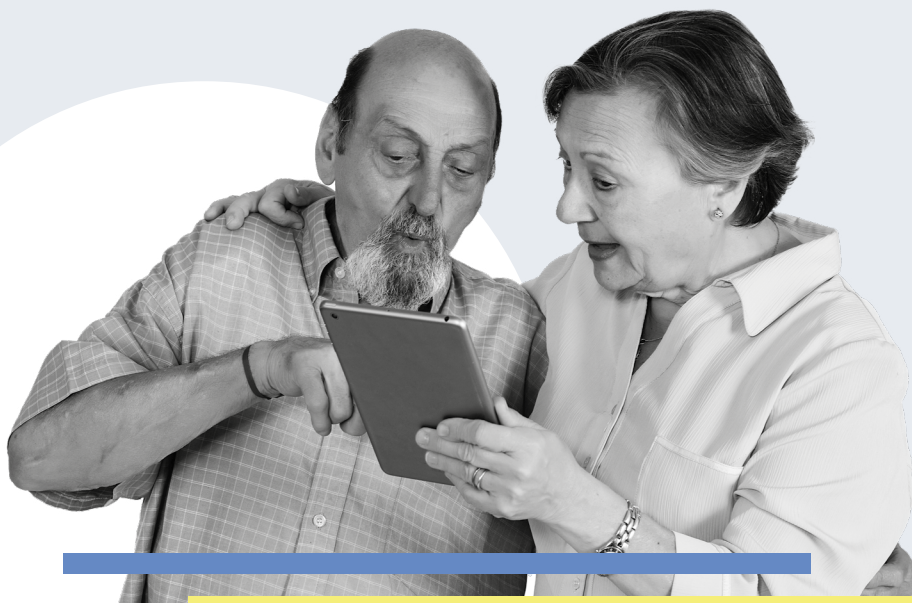
Ofiary oszustw przeżywają ogromny stres i często towarzyszy im poczucie wstydu. Wspieraj, zamiast krytykować. Porozmawiaj z taką osobą o tym, jak i gdzie zgłosić sprawę oraz podziel się radami, jak chronić się w przyszłości.



Chcesz wiedzieć więcej o bezpieczeństwie w sieci?
Sięgnij do pierwszej edycji poradnika:

#Halo! Tu cyberbezpieczny Senior!





#Halo! Tu cyberbezpieczny Senior!
Nie zadzieraj ze mną

NASK – Państwowy Instytut Badawczy
ul. Kolska 12, 01-045 Warszawa
www.nask.pl

2026

NASK



PROJEKT FINANSOWANY ZE ŚRODKÓW
MINISTERSTWA CYFRYZACJI