



Mateusz Szenawa*

**Zagrożenia gospodarcze związane
z rozwojem sztucznej inteligencji
– ze szczególnym uwzględnieniem
technologii *deepfake* –
na podstawie wybranych regulacji
prawnych Unii Europejskiej**

[Economic Threats Associated with the Development of Artificial Intelligence, with Particular Reference to Deepfake Technology, Based on Selected European Union Legal Regulations]

Abstract

This article analyses the impact of artificial intelligence, and in particular deepfake technology, on the economy. The aim of this paper is to assess whether current legal regulations at European level provide sufficient protection against the threats posed by technological development. The author presents cases in which deepfake technology can be used to extort funds or carry out other unethical activities. Key issues include the role of EU regulations – such as the Artificial Intelligence Act – in countering economic threats. Furthermore, the author highlighted the need for legislative changes that will address existing legal loopholes whilst reducing the scope of obligations imposed on businesses. It was noted that legal regulations concerning artificial intelligence are essential – whilst acknowledging that maintaining an adequate level of innovation among European businesses poses a significant challenge.

Keywords: deepfake, economy, AI Act, threats, artificial intelligence.

Wstęp

W obliczu rewolucji technologicznej, jaka dokonuje się na naszych oczach dzięki rozwojowi sztucznej inteligencji, świat staje przed nowymi zadaniami. Celem niniejszego artykułu jest dokonanie analizy skuteczności regulacji unijnych, a w konsekwencji również uregulowań krajowych w zakresie ochrony

* **Mateusz Szenawa** – prawnik i matematyk, asystent w Instytucie Nauk Prawnych Uniwersytetu Opolskiego (afiliacja); <https://orcid.org/0009-0005-3393-1316>; mateusz.szenawa@uni.opole.pl / lawyer and mathematician, research assistant at the Institute of Legal Studies, University of Opole (affiliation).

prawnej przed gospodarczymi skutkami wykorzystania technologii deepfake. Praca odpowiada na główne pytanie badawcze: czy i w jakim zakresie obowiązujące mechanizmy prawne gwarantują adekwatną, wystarczającą ochronę osób fizycznych oraz podmiotów gospodarczych przed negatywnymi skutkami wykorzystania narzędzi sztucznej inteligencji umożliwiających tworzenie nieprawdziwych treści przez: oszustwa, dezinformację, naruszenie czci czy szkodenie reputacji. Naturalną kontynuacją rozważań nad głównym problemem badawczym jest rozpatrzenie, w jakim stopniu brak kompleksowej – a zatem pozbawionej luk – definicji pojęcia deepfake utrudnia ochronę praw jednostki oraz wpływa na zdolność systemu prawnego do przeciwdziałania gospodarczym zagrożeniom, rozumianym jako zestaw ryzyk ekonomicznych oraz prawnych wynikających z nieetycznego lub przestępczego wykorzystania narzędzi sztucznej inteligencji. W konsekwencji: należy rozpoznać, jakie mechanizmy prawne powinny funkcjonować, aby zapobiec negatywnym skutkom wykorzystywania tej technologii oraz jak ich wprowadzenie wpływa na innowacyjność i konkurencyjność przedsiębiorstw prowadzących działalność na terenie Unii Europejskiej. To właśnie powyższe pytania stanowią punkt wyjścia dla dalszych rozważań.

W artykule dokonano analizy poziomu ochrony prawnej przed cyfrowymi zagrożeniami wynikającymi z rozwoju AI, ze szczególnym uwzględnieniem regulacji unijnych w zakresie ochrony przed oszustwami, manipulacjami rynkowymi oraz naruszeniem wizerunku. Jednocześnie dynamika rozwoju nowych technologii stawia przed ustawodawcą wyzwanie utrzymania normatywnej adekwatności przepisów prawa. W ramach analiz dokonano rozważań *de lege ferenda*, identyfikując istotną, dotąd nieopisaną lukę definicyjną występującą w przepisach unijnych. Postulowane zmiany mają na celu wzmocnienie bezpieczeństwa gospodarczego w sytuacjach zderzenia się przedsiębiorców i ich pracowników z nowymi technologiami, w tym z zagrożeniami z nich płynącymi.

W opracowaniu jako podstawową zastosowano metodę dogmatycznoprawną – analizującą aktualne regulacje unijne w świetle ich zgodności z celem, jakim jest ochrona uczestników rynku. Metodę dogmatycznoprawną uzupełniono metodą historycznoprawną – pozwalającą uchwycić ewolucję przepisów w czasie – oraz elementami analizy porównawczej, w ramach której zestawiono wybrane rozwiązania unijne z elementami regulacji obowiązujących w USA i Wielkiej Brytanii. W wybranych szczegółowych kwestiach przywołano też kierunki regulacyjne obowiązujące w Republice Korei (państwie o dalece zaawansowanym systemie prawnym dotyczącym technologii w regionie Azji). W pierwszej kolejności zwrócono uwagę na problematykę zdefiniowania pojęcia deepfake i sztucznej inteligencji oraz na różnorodne aspekty wykorzystania deepfake'ów – od dezinformacji, przez oszustwa finansowe, po potencjalne skandale związane z naruszaniem dóbr osobistych wpływowych

osób. Rozwój technologii deepfake stawia pytania o granice odpowiedzialności prawnej za treści generowane przez systemy AI, rodząc potrzebę rewizji obowiązujących obecnie w tym zakresie przepisów prawa. Ostatnia część artykułu poświęcona została krytycznej analizie obowiązujących regulacji prawnych pod kątem ich adekwatności wobec zidentyfikowanych wyżej zagrożeń.

Pojęcie deepfake w kontekście rozwoju sztucznej inteligencji

Pojęcie deepfake powstało z połączenia dwóch odrębnych wyrazów „deep learning” – czyli: głębokie uczenie – oraz „fake”, co oznacza: fałszywy. Deep learning to przetwarzanie danych za pomocą sztucznych sieci neuronowych. „Modele uczenia głębokiego pobierają informacje z wielu źródeł i analizują te dane w czasie rzeczywistym bez konieczności interwencji ze strony człowieka”¹. Z tego wynika, że technika ta opiera się na sztucznej inteligencji, a konkretnie uczeniu maszynowym². W ogólności zatem deepfake polega na wytworzeniu fałszywego – to znaczy takiego, który w rzeczywistym świecie nie wystąpił – obrazu lub treści wideo z wykorzystaniem narzędzi związanych ze sztuczną inteligencją. Najczęściej stworzenie takiej treści polega na zastąpieniu w istniejącym materiale jednej osoby inną lub stworzeniu całkowicie oryginalnych i nowych treści, w których ludzie zachowują się tak, jak tego nigdy w rzeczywistości nie robili³. Dzieje się to z wykorzystaniem funkcjonalności konkretnego komponentu sztucznej inteligencji, a dokładnie deep learningu, który jest jednocześnie podzbiorem machine learningu, czyli wspomnianego uczenia maszynowego. W ramach powyższych rozważań istotne staje się również rozróżnienie pojęcia deepfake od fake newsów. Fake news stanowi przekaz medialny o charakterze dezinformującym, który często zawiera prawdziwe elementy, lecz jego celem jest manipulacja faktami⁴. Deepfake natomiast w głównej mierze polega na wytworzeniu nieprawdziwych obrazów lub treści wideo.

Należy zaznaczyć, że technologia deepfake może zostać użyta w dwojakim kontekście – pozytywnym i negatywnym. Do podobnych wniosków doszła I. Dąbrowska, która wyróżniła cztery pierwotne intencje, dla jakich deepfake

¹ A. Piecuch, *Sztuczna inteligencja w perspektywie społecznej*, „Edukacja Ustawiczna Dorosłych” 2023, 123, 4, ss. 13–25.

² J. Brandon, *Terrifying high-tech porn: Creepy ‘deepfake’ videos are on the rise*, <https://www.foxnews.com/tech/terrifying-high-tech-porn-creepy-deepfake-videos-are-on-the-rise> [dostęp: 7.12.2024].

³ D. Kumari, *Deepfake Technology and Legal Issues*, ‘LawFoyer International Journal of Doctrinal Research’ 2024, 2, 1, s. 235.

⁴ K. Kiepiński, *Deepfake jako narzędzie do przekazywania informacji fałszywej i domniemanej. Analiza prawnokarna i cybernetyczna*, „Kwartalnik Krajowej Szkoły Sądownictwa i Prokuratury” 2023, 3, 51, s. 89.

został stworzony: rozrywkę, edukację, dezinformację i dyskredytację⁵. Jednocześnie należy zaznaczyć, że przywołany katalog nie jest wyczerpujący. Używanie tej technologii rzeczywiście następuje w świecie rozrywki, filmowym czy edukacyjnym (kontekst pozytywny). Ponadto *deepfake* może zostać wykorzystany w celach dezinformacyjnych, mających wprowadzać w błąd – z posłużeniem się na przykład metodami socjotechnicznymi, wpływającymi na proces decyzyjny konkretnych jednostek; może również być użyty w celu dokonania cyberprzestępstw (kontekst negatywny). Technologia *deepfake* staje się przydatnym narzędziem w atakach mających na celu wyłudzenie pieniędzy lub innej wartości, np. danych osobowych czy informacji objętych tajemnicą przedsiębiorstwa. Pomimo to *deepfake* nie został zdefiniowany ani w polskim prawie karnym, ani cywilnym. W ramach kodeksu karnego oraz kodeksu cywilnego nie został też do tej pory wyodrębniony żaden przepis odnoszący się bezpośrednio do czynów związanych z wykorzystaniem jednego z narzędzi sztucznej inteligencji, jakim jest technologia *deepfake*.

Dopiero przepisy na szczelbu europejskim – a dokładnie rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 z 13 czerwca 2024 r. w sprawie ustanowienia zharmonizowanych przepisów dotyczących sztucznej inteligencji (dalej: AI Act lub rozporządzenie 2024/1689) – wprowadza w art. 3 pkt 60 definicję legalną tego pojęcia; według niej „*deepfake*” oznacza wygenerowane przez AI lub zmanipulowane przez AI obrazy, treści dźwiękowe lub treści wideo przypominające istniejące osoby, przedmioty, miejsca, podmioty lub zdarzenia, które odbiorca mógłby niesłusznie uznać za autentyczne lub prawdziwe⁶. Definicja ta zmieniała się w trakcie prac nad AI Act. Przykładem może być chociażby proponowana treść art. 52 ust. 3 wniosku Komisji Europejskiej w sprawie AI Act z 21 kwietnia 2021 r., gdzie *deepfake* został zdefiniowany jako system sztucznej inteligencji, który generuje obrazy, treści dźwiękowe lub treści wideo łudząco przypominające istniejące osoby, obiekty, miejsca lub inne podmioty lub zdarzenia lub który tymi obrazami i treściami manipuluje – przez co osoba będąca ich odbiorcą mogłaby niesłusznie uznać je za autentyczne lub prawdziwe⁷. Łatwo jednak zauważyć, że w obu tych definicjach możemy znaleźć hasła przewodnie, które są podstawą właściwego zdefiniowania, czym jest *deepfake*. Aby daną treść uznać za materiał *deepfake*, muszą zostać spełnione następujące przesłanki:

⁵ I. Dąbrowska, *Deepfake – nowy wymiar internetowej manipulacji*, „Zarządzanie Mediami” 2020, 8, 2, s. 96.

⁶ Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 z 13 czerwca 2024 r. w sprawie ustanowienia zharmonizowanych przepisów dotyczących sztucznej inteligencji oraz zmiany rozporządzeń (WE) nr 300/2008, (UE) nr 167/2013, (UE) nr 168/2013, (UE) 2018/858, (UE) 2018/1139 i (UE) 2019/2144 oraz dyrektyw 2014/90/UE, (UE) 2016/797 i (UE) 2020/1828 (akt w sprawie sztucznej inteligencji), (Dz.U. L 2024/1689, 12.07.2024) zwanego dalej Rozporządzeniem 2024/1689.

⁷ Wniosek rozporządzenie Parlamentu Europejskiego i Rady ustanawiające zharmonizowane przepisy dotyczące sztucznej inteligencji (akt w sprawie sztucznej inteligencji) i zmieniające niektóre akty ustawodawcze unii, (COM/2021/206 final)

- a) stworzenie lub zmanipulowanie przez AI;
- b) ukazanie treści łudząco zbliżonych do elementów rzeczywistych lub do istniejących osób;
- c) istnienie możliwości niesłusznego uznania treści za prawdziwą.

Jako że materiały *deepfake* powstają z wykorzystaniem sztucznej inteligencji, należy się zastanowić, czym sztuczna inteligencja naprawdę jest. Pojęcie sztucznej inteligencji (*artificial intelligence*, AI) zostało zdefiniowane na wiele różnych sposobów. Niemniej dla rozważań wynikających z niniejszego artykułu wystarczająca będzie definicja wynikająca z analiz brytyjskiego matematyka Alana Turinga. Zgodnie z jego koncepcją „sztuczna inteligencja to zdolność maszyny do naśladowania lub imitowania ludzkiej inteligencji”⁸. Wynika ona wprost ze słynnego testu Turinga, „który ma definiować normy dla maszyny odznaczającej się inteligencją”⁹. Taka definicja wydaje się uniwersalna i – co wręcz niespotykane w dziedzinie nowych technologii – ponadczasowa (wynika bowiem z pracy „*Computing Machinery and Intelligence*”, opublikowanej w 1950 roku)¹⁰. Inne próby zdefiniowania sztucznej inteligencji wielokrotnie kończą się stworzeniem skomplikowanych struktur logicznych, które zamiast prowadzić do ujednolicenia definicji, wprowadzają niepotrzebne zamieszanie w debacie publicznej¹¹. Zdarza się również, że do zdefiniowania sztucznej inteligencji używana bywa ta autorstwa Johna McCarthy’ego, zaproponowana w 1955 roku, która uważana jest za jedną z pierwszych, a przez J. Widłę za pierwszą definicję sztucznej inteligencji¹². Ponieważ jednak John McCarthy w swojej propozycji był zgodny z rozważaniami A. Turinga – z których wynika, że maszyna jest inteligentna, gdy człowiek nie jest w stanie odróżnić, czy udzielona odpowiedź na zadane pytanie pochodzi od człowieka, czy tej maszyny – to trudno uznać, że definicja ta jest pierwsza w sensie koncepcyjnym. Nawet jeśli A. Turing nie wprowadził formalnego sensu *stricto* pojęcia sztucznej inteligencji, to przyczynił się do zdefiniowania kryteriów inteligencji maszyn¹³.

⁸ T. Zalewski, *Definicja sztucznej inteligencji* [w:] L. Lai, M. Świerczyński (red.), *Prawo sztucznej inteligencji*, Warszawa 2020, ss. 5–9.

⁹ I. Stewart, *Significant Figures Lives and Works of Trailblazing Mathematics*, New York 2017, s. 325 [wyd. pol. Krótka historia wielkich umysłów. Genialni matematycy i ich arcydzieła, tłum. U. Seweryńska i M. Seweryński, Warszawa 2019].

¹⁰ *Ibid.*

¹¹ Por. T. Zalewski, *Definicja...*, *ibid.*; D. Spałek, J. Rzymowski, *Jak nie stosować AI Act na podstawie definicji systemu sztucznej inteligencji*, „*Prawo Nowych Technologii*” 2024, 2, s. 21.

¹² J. Widło, *Odpowiedzialność za szkodę wyrządzoną przez pojazdy autonomiczne w świetle zmian prawa UE*, „*Forum Prawnicze*” 2025, 2, s. 7.

¹³ A. Turing, *Computing Machinery and Intelligence*, „*Mind*” 1950, 59, 236, s. 433–435.

Zagrożenia wynikające z technologii deepfake w kontekście gospodarczym

Dezinformacja

Badacze zwracają uwagę na zagrożenia związane z dezinformacją w przestrzeni publicznej z wykorzystaniem fake newsów czy materiałów wygenerowanych za pomocą technologii deepfake. Niewątpliwie „rozwój nowych technologii (...) może zwiększyć zagrożenie dla stabilności politycznej i społecznej”¹⁴. W tym kontekście aspekt gospodarczy pozostaje niedostatecznie zbadany. Niniejsze opracowanie uzupełnia tę lukę poprzez analizę zagrożeń dotyczących przedsiębiorców, ponieważ dezinformację znaną ze świata polityki można przełożyć na świat biznesu. Przedsiębiorcy mogą tworzyć fałszywie pozytywne informacje o swojej firmie, realizując na przykład nieprawdziwy, pozytywny marketing, jak również czarny PR konkurencji, gdzie materiały deepfake mogą być używane przez konkurencyjne firmy przygotowujące kompromitujące lub dezinformujące materiały w sprawie wyników finansowych, polityki firmy czy sposobu traktowania przez tę firmę pracowników. Podkreślenia wymaga fakt, że nie tylko przedsiębiorcy mogą być twórcami tego typu treści, ale również osoby fizyczne (pracownicy, inwestorzy, osoby o nieprzychylnym usposobieniu w stosunku do konkretnej firmy), które mogą mieć w tym interes osobisty. Temat ten jest pobieżnie traktowany przez przepisy prawne, ponieważ nie dotyczy ogółu społeczeństwa, porządku publicznego czy bezpieczeństwa państwa, a miałby chronić interes konkretnej jednostki. Tymczasem im większe przedsiębiorstwo stanie się ofiarą albo będzie twórcą ataku deepfake, tym reperkusje społeczne mogą być większe. Zmasowany i skrupulatnie przygotowany dezinformacyjny atak deepfake może wpłynąć znacząco na kurs akcji spółek kapitałowych, w tym spółek strategicznych dla bezpieczeństwa państwa.

Również biała księga w sprawie sztucznej inteligencji zwraca uwagę na wyzwania stawiane przed nami przez rozwój sztucznej inteligencji, jednak w niewystarczającym stopniu identyfikuje ona potencjalne zagrożenia gospodarcze, które mogą wystąpić w sytuacji przestępczego wykorzystania nowoczesnych narzędzi technologicznych¹⁵. Komisja Europejska zauważyła, że narzędzia w zakresie sztucznej inteligencji mogą przyczynić się do lepszej

¹⁴ D. Mierzejewski, J. Nawrotkiewicz, W obliczu wspólnej polityki UE wobec Chin: „mission impossible”?, „Pulaski Policy Papers” 2024, XI, s. 20.

¹⁵ Biała księga w sprawie sztucznej inteligencji. Europejskie podejście do doskonałości i zaufania (COM/2020/65 final/2), zwana dalej białą księgą.

ochrony obywateli UE przed przestępczością i aktami terrorystycznymi¹⁶. Nie zmienia to faktu, że spośród katalogu przykładowo wymienionych czynności, w których sztuczna inteligencja może okazać się przydatna, brakuje informacji o zapobieganiu i wykrywaniu dezinformacji dokonanej przez tworzenie fałszywych treści *deepfake*. Pozytywnie należy jednak ocenić fakt zauważenia przez Komisję Europejską, że mogą być konieczne dalsze ustalenia w celu zapobiegania i zwalczania stosowania AI w celach przestępczych¹⁷. Również część białej księgi dotycząca bezpośrednio zagrożeń dla bezpieczeństwa wymienia te zagrożenia, które „mogą być spowodowane wadami w projekcie technologii sztucznej inteligencji, być związane z problemami z dostępnością i jakością danych lub z innymi problemami wynikającymi z uczenia się maszyn”¹⁸. Należy zaznaczyć, że biała księga miała za zadanie przedstawić warianty strategiczne dotyczące sposobów osiągnięcia dwóch celów, którymi są promowanie sztucznej inteligencji oraz zajęcie się zagrożeniami wynikającymi z konkretnych jej zastosowań¹⁹.

Potencjalne skompromitowanie wpływowych osób

Naturalną kontynuacją tworzenia dezinformujących treści *deepfake* jest naruszenie dóbr osobistych poszczególnych osób. W poniższych rozważaniach przeanalizowano potencjalne konsekwencje prób skompromitowania wpływowych osób ze świata biznesu. Według danych zebranych przez NASK grupą społeczną najbardziej narażoną na tworzenie materiałów *deepfake* są dziennikarze – 32%, następnie politycy – 21%. Trzecią grupą społeczną są biznesmeni – 11%²⁰. Wynik ten pokazuje skalę zagrożenia, do której zarówno przepisy prawa krajowego, jak i unijnego nie są przystosowane. Jak wskazuje A. Ziobroń, w ramach przepisów krajowych przewidziano ochronę została przewidziana w art. 212 oraz art. 216 k.k.²¹ Ponadto wizerunek i cześć są chronione na podstawie art. 23 k.c., co daje ofercie kumulatywną ochronę prawną. Ten obecny stan prawny należy wszakże uznać za nieadekwatny do

¹⁶ Biała księga zawiera informację, że sztuczna inteligencja „mogłyby na przykład pomóc w identyfikowaniu internetowej propagandy terrorystycznej, wykrywaniu podejrzanych transakcji sprzedaży niebezpiecznych produktów, identyfikowaniu niebezpiecznych ukrytych przedmiotów lub nielegalnych substancji lub produktów czy też oferować pomoc obywatelom w sytuacjach nadzwyczajnych i wspomagać służby interwencyjne” (s. 2).

¹⁷ Biała..., s. 3.

¹⁸ Biała..., s. 14.

¹⁹ A. Zielińska, *Korzystanie z treści chronionych a sztuczna inteligencja – jak daleko sięga odpowiedzialność dostawcy usług udostępniania treści online w świetle art. 17 dyrektywy 2019/790?*, „Studia de Cultura” 2022, 14, 2, s. 136.

²⁰ NASK ostrzega: przestępcy tworzą coraz bardziej pomysłowe oszustwa *deepfake*, <https://nask.pl/aktualnosci/nask-ostrega-przestepcy-tworza-coraz-bardziej-pomyslowe-oszustwa-deepfake/>, [dostęp: 12.12.2024].

²¹ A. Ziobroń, *Deepfake a prawo karne. Uwagi „de lege lata” i „de lege ferenda” dotyczące fałszywej pornografii*, „Studenckie Prace Prawnicze, Administratywistyczne i Ekonomiczne” 2021, 37, s. 225.

skali zagrożenia, ochrona prawna w dużej mierze jest bowiem tylko pozorna, gdyż ze względu na charakter czynu znalezienie sprawcy jest utrudnione, a potencjalne konsekwencje karne pozostają niewspółmierne do skali społecznej szkodliwości. Do podobnych wniosków dociera też G.G. Vera – twierdząc, że „ochrona ta jest fragmentaryczna i niewystarczająca”²².

Podobnie należy również spojrzeć na unijne regulacje (AI Act), które zwracają uwagę na fakt, że sztuczna inteligencja stwarza nowe rodzaje ryzyka – polegające na podawaniu informacji wprowadzających w błąd i na manipulacji na dużą skalę, oszustwach, podszywaniu się pod inne osoby i wprowadzaniu w błąd konsumentów²³. W ramach tych zagrożeń AI Act zakłada konieczność zobowiązania dostawców systemów do wbudowania takich narzędzi, które pozwolą wykrywać, że poszczególne treści zostały wygenerowane przez sztuczną inteligencję, a nie przez człowieka. AI Act w zdecydowanie większym stopniu zawiera opis zagrożeń związanych z dezinformacją, podszywaniem się pod inną osobę i tworzeniem treści deepfake, niż przewidywała to przywołana wyżej biała księga. W przypadku tworzenia fałszywych treści z wykorzystaniem technologii deepfake należy zwłaszcza wyróżnić te szczególnie szkodliwe, bo atakujące bezpośrednio osobę w celu jej zdyskredytowania – np. przez stworzenie materiału typu deepfake porno, gdzie znana osobistość „występuje” w spreparowanym za pomocą sztucznej inteligencji materiale pornograficznym (główną ofiarą takich materiałów są kobiety), czy też przez wygenerowanie fałszywych nagrań zawierających treści, z których wynikałyby przykładowo rzekome obyczajowe dewiacje lub niewłaściwe traktowanie podwładnych przez właściciela firmy. Niestety możliwości wytworzenia takich treści ograniczone są wyłącznie pomysłowością twórcy takich materiałów. Jeszcze innym zagrożeniem, które należy odróżnić od powyższych, jest wykorzystanie wizerunku znanych przedsiębiorców czy innych osób ze świata biznesu w celu dokonania oszustw finansowych i wyłudzeń pieniędzy oraz danych.

Cyberoszustwa

Próby wyłudzenia środków finansowych z wykorzystaniem nowoczesnych technologii należy podzielić na dwie podstawowe kategorie. Pierwszą jest sytuacja, gdy ofiarą staje się konkretna jednostka. Dotyczy to przede wszystkim sytuacji, w których cyberoszuści – tworząc spreparowane treści deepfake – zachęcają do fałszywych inwestycji z wykorzystaniem wizerunku znanych przedsiębiorców. Tylko w pierwszej połowie 2024 roku NASK wyszczególnił

²² G.G. Vera, Deep fake (!) – postęp technologiczny a prawo karne, „Zeszyty Naukowe Uniwersytetu Rzeszowskiego – Seria Prawnicza” 2024, 1, 44, s. 50.

²³ Motyw 133 preambuły rozporządzenia 2024/1689.

jedenastu polskich przedsiębiorców, którzy padli ofiarą *deepfake*²⁴. Tego typu działania – oprócz strat finansowych poniesionych przez ofiary cyberoszustów – przyczyniają się również do utraty zaufania do poszczególnych przedsiębiorców. Mimo że jest jasne, iż tego typu fałszywe zachęty inwestycyjne nie są firmowane przez przedsiębiorców, których wizerunek wykorzystano, wystarczy, że chociażby część społeczeństwa nie będzie tego świadoma i postanowi rozpowszechniać nieprawdziwe informacje dotyczące danej osoby (osób). Pochodną tego typu zachowań może stać się utrata zaufania do konkretnej firmy, a to może bezpośrednio wpływać na wyniki finansowe przedsiębiorstwa.

Drugą kategorią cyberoszustw, na które należy zwrócić uwagę, są te, w których szkoda dotyczy bezpośrednio przedsiębiorstwa. Ofiarą nadal jest konkretna jednostka, ale środki finansowe lub dane, które cyberoszuści chcą w tej sytuacji wyłudzić, nie są jej własnością. Pracownik, który padł ofiarą cyberprzestępców, może w ramach swoich obowiązków zawodowych jedynie dysponować majątkiem firmowym lub danymi stanowiącymi tajemnicę przedsiębiorstwa. Jako przykład tego typu działania należy podać przypadek oszustwa z użyciem technologii *deepfake*, w którym brytyjska firma projektowa Arup straciła 25 mln USD²⁵. Pracownik sądził, że uczestniczy w poufnej transakcji z dyrektorem i innymi pracownikami firmy²⁶. Omówiony przypadek jest szczególnie wymowny z perspektywy regulacji obowiązujących w Zjednoczonym Królestwie. Arup jest podmiotem brytyjskim z główną siedzibą w Londynie. Wielka Brytania przyjęła w ramach swojego ustawodawstwa Online Safety Act 2023, mający na celu nałożenie obowiązków ochronnych na platformy internetowe w zakresie nielegalnych treści udostępnionych przez użytkowników²⁷. Jednak przyjęty model dotyczy tylko reaktywnej moderacji treści i znajduje zastosowanie wyłącznie wtedy, gdy dany materiał będzie już dostępny online²⁸. W kontekście oszustw finansowych taki charakter regulacji okazuje się niewystarczający, ponieważ szkoda materializuje się w momencie dokonywania przelewu przez ofiarę, a nie w momencie udostępnienia treści *deepfake* w sieci²⁹. Wątpliwe jest przy tym, czy tak wygenerowana treść

²⁴ Lista osób, których wizerunek przestępcy wykorzystali w oszustwach *deepfake*, <https://nask.pl/aktualnosci/nask-ostrzega-przestepcy-tworza-coraz-bardziej-pomyslowne-oszustwa-deepfake/> [dostęp: 12.12.2024].

²⁵ D. Jakubczak, *Deepfake. Oszustwo na sklonowanego dyrektora*, <https://www.dw.com/pl/deepfake-oszustwo-na-sklonowanego-dyrektora/a-69116765> [dostęp: 16.12.2024].

²⁶ Oszuści wykorzystali cyfrowo sklonowany wizerunek i głos dyrektora finansowego firmy, by podczas fałszywej wideokonferencji przekonać pracownika w Hongkongu do wykonania piętnastu przelewów na różne konta bankowe.

²⁷ Online Safety Act 2023, c. 50, Zjednoczone Królestwo, <https://www.legislation.gov.uk/ukpga/2023/50/contents/enacted> [dostęp: 10.05.2026].

²⁸ B. Kira, *When Non-Consensual Intimate Deepfakes Go Viral: The Insufficiency of the UK Online Safety Act*, 'Computer Law & Security Review' 2024, 54, art. 106024, s. 7.

²⁹ Opisany akt prawny daje możliwość ochrony dzieci przed dostępem do szkodliwych treści oraz zmierza do ograniczenia rozpowszechniania intymnych treści *deepfake*. Trafnie wskazuje się jednak, że regulacja nie zawiera wystarczających mechanizmów prewencyjnych i obejmuje wyłącznie ułamek rzeczywistych zagrożeń występujących w środowisku online – zob. B. Kira, *When Non-Consensual...*, ibid.

w ogóle zostałyby przez cyberoszustów udostępniona. Ponadto przy atakach na międzynarodowe podmioty powstaje problem reguł kolizyjnychprawnych i stosowania prawa właściwego; choć bowiem straty dotyczą podmiotu brytyjskiego, przelewy trafiły na konta bankowe znajdujące się w Hongkongu. Podkreślenia wymaga, że przedsiębiorcy – a w szczególności korporacje – mają wypracowane wewnętrzne systemy i procedury mające zapobiegać cyberoszustwom. Atak na przedsiębiorstwo wymaga zdecydowanie większej wprawy i profesjonalizmu niż ten, w którym ofiarą miałyby być osoba fizyczna. Jednocześnie nie ulega wątpliwości, że obecne możliwości sztucznej inteligencji są już wystarczające, aby takich oszustw dokonać.

Cyberoszustwa z wykorzystaniem technologii deepfake stanowiły dotychczas wyłącznie wycinek ogółu cyberprzestępstw. Niemniej przywołane poniżej dane uzasadniają potraktowanie tego zjawiska jako priorytetowego obszaru do uregulowania. W 2022 roku w USA zidentyfikowano 300 497 ofiar phishingu, a wartość skradzionych środków finansowych oszacowano na 52 mln USD³⁰. Powyższe potwierdza również raport FBI, który wskazuje, że dokładna wartość utraconych środków w ramach phishingu wynosiła w 2022 roku 52 089 159 USD³¹. Zatem średnia wysokość utraconej przez jedną ofiarę kwoty to wartość około 173 USD. Zestawiając to z faktem, że przy wykorzystaniu zaawansowanych narzędzi deepfake jedna osoba była skłonna dokonać łącznych przelewów na kwotę 25 mln USD – co stanowi 48% łącznej kwoty, jaką wszystkie zidentyfikowane ofiary phishingu w Stanach Zjednoczonych straciły przez cały 2022 rok – widoczna staje się skala zagrożenia.

Kolejne raporty – analizując wyłącznie dane od 2023 roku, które są agregowane w jednolity sposób – wskazują jednocześnie na spadek liczby zgłaszanych przestępstw dokonanych w ramach phishingu/spoofingu przy jednoczesnym wzroście wartości utraconych przez ofiary środków finansowych³². Należy zaznaczyć, że raport z 2022 roku osobno podaje wysokości utraconych kwot przez ofiary phishingu i spoofingu, natomiast w raporcie z 2024 r. oraz 2025 r. podane wartości dotyczą zsumowanych strat finansowych ofiar phishingu i spoofingu³³. Istotne jest również, że w 2023 roku w Stanach Zjednoczonych łączne straty ofiar tych dwóch procedurów wynosiły 18 728 550 USD, co stanowi ok. 6–7 mln USD mniej niż wyłącznie jeden atak deepfake na przedsiębiorstwo³⁴. Oszustwa z wykorzystaniem wideokonferencji wygenerowanych

³⁰ M. Bukowski, Zwalczanie cyberprzestępczości ekonomicznej przy wykorzystaniu sztucznej inteligencji (AI), „Przegląd Policyjny” 2023, 151, 3, s. 163.

³¹ FBI, Internet Crime Report 2022, https://www.ic3.gov/AnnualReport/Reports/2022_ic3report.pdf, s. 22 [dostęp: 24.10.2025].

³² FBI, Internet Crime Report 2024, https://www.ic3.gov/AnnualReport/Reports/2024_IC3Report.pdf, s. 18 [dostęp: 24.10.2025]; FBI, Internet Crime Report 2025, https://www.ic3.gov/AnnualReport/Reports/2025_IC3Report.pdf, s. 25 [dostęp: 8.05.2026].

³³ Ibid.

³⁴ Ibid.

przez sztuczną inteligencję są zdecydowanie trudniejsze do przygotowania przez oszustów, ale zarazem potencjalne straty finansowe ich ofiar są niewspółmiernie wyższe.

Inne cyberprzestępstwa

Sztuczna inteligencja może być wykorzystywana przez cyberprzestępców na wiele jeszcze innych sposobów. Materiały stworzone za pomocą technologii deepfake mogą zostać wykorzystane w celu zwiększenia skuteczności ataków phishingowych. Ponadto sztuczna inteligencja może zostać wykorzystana do zautomatyzowania tworzenia nieprawdziwych informacji/recenzji na temat produktów lub usług świadczonych przez konkretne przedsiębiorstwo. Dodatkowo takie opinie mogą być wzmacniane przez stworzenie fałszywego wideo, w którym wygenerowana osoba pozytywnie recenzuje wybrane przedsiębiorstwo lub negatywnie odnosi się do działań firm konkurencyjnych. Należy również zauważyć całą rzeszę możliwości AI poza technologią deepfake, które umożliwiają jeszcze skuteczniejsze popełnianie cyberprzestępstw. Wśród nich można wymienić między innymi generowanie fałszywych dokumentów, możliwość manipulowania rynkiem kapitałowym przez tworzenie fałszywych informacji czy też wykorzystywanie botów do wyłudzeń danych. Jednak z racji obszerności tematu przestępstwa te winny zostać rozważone w ramach odrębnej pracy badawczej.

W tym miejscu należałoby zwrócić uwagę na zjawisko Crime/Cybercrime as a Service, które polega na „sprzedaży dostępu do narzędzi i wiedzy potrzebnej do przeprowadzenia ataków”³⁵. W rezultacie osoby z ograniczoną wiedzą i możliwościami mogą przeprowadzić ataki ze względną łatwością³⁶. Tego typu usługa może również polegać na odprzedaży właściwych narzędzi do wytworzenia w prosty sposób dowolnych treści deepfake. Możliwość ta generuje ryzyka gospodarcze związane z chęcią między innymi zaszkodzenia konkretnej firmie czy konkretnej osobie poprzez wytworzenie szkodliwych materiałów przez osobę niemającą wystarczających umiejętności, aby w pełni samodzielnie tak sfałszowany materiał deepfake przygotować. Poznawszy już główne zagrożenia gospodarcze, należy się zastanowić, w jakim stopniu obecne przepisy prawne są na nie przygotowane.

³⁵ Crime as a service (CaaS), <https://www.comcert.pl/slownik/crime-as-a-service-caas/> [dostęp: 17.12.2024].

³⁶ A. Unni, What You Need To Know About Crime As A Service (CSaaS), <https://www.stickmancyber.com/cybersecurity-blog/what-you-need-to-know-about-crime-as-a-service-csaas> [dostęp: 17.12.2024].

Akt o sztucznej inteligencji jako próba odpowiedzi na zidentyfikowane zagrożenia

Przepisy AI Act należy rozważyć w kontekście zagrożeń opisanych powyżej. Komisja Europejska zwróciła uwagę na ryzyko dezinformacji i manipulacji, które mogą być szerzone za pomocą sztucznej inteligencji, w tym technologii deepfake. AI Act zwraca uwagę, że sztuczna inteligencja może być wykorzystywana niewłaściwie i może dostarczać nowych, potężnych narzędzi do praktyk manipulacji, wyzyskiwania i kontroli społecznej³⁷. Jednocześnie rozporządzenie wprost zakłada, że praktyki takie powinny być zakazane jako niezgodne z unijnymi wartościami. W każdym razie nie ma znaczenia, czy dostawca lub podmiot stosujący mieli zamiar wyrządzić poważną szkodę; istotny jest fakt, że szkoda wynika z praktyk manipulacyjnych opartych na AI lub z ich wykorzystywania. Zakazy dotyczące takich praktyk w zakresie AI stanowią uzupełnienie przepisów zawartych w dyrektywie Parlamentu Europejskiego i Rady 2005/29/WE dotyczącej nieuczciwych praktyk handlowych stosowanych przez przedsiębiorstwa wobec konsumentów na rynku wewnętrznym³⁸, w szczególności przepisu zakazującego stosowania we wszelkich okolicznościach nieuczciwych praktyk handlowych powodujących dla konsumentów szkody ekonomiczne lub finansowe – niezależnie od tego, czy praktyki te stosuje się za pomocą systemów AI czy w innym kontekście³⁹. Na szczęblu unijnym zwrócono zatem uwagę na zagrożenia gospodarcze i możliwość powstania szkód materialnych oraz finansowych przy niewłaściwym stosowaniu narzędzi AI.

Pozytywnie należy ocenić fakt, że AI Act zawiera neutralne technologicznie definiowanie pojęć, dzięki czemu uzyskuje szerokie zastosowanie oraz wieloletnią aktualność przyjętych pojęć. W rozporządzeniu 2024/1689 położono mocniejszy nacisk na opis działania konkretnych narzędzi niż na ich enumeratywne wymienianie w ramach zamkniętego katalogu – co powinno tym samym wydłużyć aktualność przepisów AI Act. Podobnie zauważa też Dorota Czyżewska-Misztal, stwierdzając, że projekt AI Act zapewnia „horyzontalne ramy prawne dotyczące sztucznej inteligencji, które mają zapewnić pewność prawa”⁴⁰. Pomimo że końcowa treść AI Act różni się znacząco od projektu, udało się zachować takie definiowanie pojęć, które nie powinno ulec szybkiej dezaktualizacji.

Należy również zaznaczyć bardzo ważną w kontekście technologii deepfake zmianę treści rozporządzenia 2024/1689 względem treści, która wy-

³⁷ Motyw 28 preambuły rozporządzenia 2024/1689.

³⁸ Dz. Urz. UE 2005 L 149/22.

³⁹ Motyw 29 preambuły rozporządzenia 2024/1689.

⁴⁰ D. Czyżewska-Misztal, Europejskie podejście do sztucznej inteligencji [w:] Skutki i zagrożenia cywilizacji informacyjnej, K.A. Nawrot, K. Prandacki (red.), Komitet Prognoz PAN 2023, s. 123.

stąpiła w jego projekcie. Załącznik III do projektu AI Act wymieniał systemy sztucznej inteligencji przeznaczone do wykorzystania przez organy ścigania do wykrywania treści stworzonych z wykorzystaniem technologii *deepfake*, o których mowa w art. 52 ust. 3, jako systemy wysokiego ryzyka⁴¹. Jednocześnie art. 52 ust. 3 projektu zakładał, że użytkownicy systemu sztucznej inteligencji, który generuje obrazy, treści dźwiękowe lub treści wideo łudząco przypominające istniejące osoby, obiekty, miejsca lub inne podmioty lub zdarzenia lub który tymi obrazami i treściami manipuluje – przez co osoba będąca ich odbiorcą mogłaby niesłusznie uznać je za autentyczne lub prawdziwe („*deepfake*”), ujawnią, że dane treści zostały wygenerowane lub zmanipulowane przez system sztucznej inteligencji⁴². Sytuacja ta zakładała, że systemy służące do wykrywania przestępstw z wykorzystaniem technologii *deepfake* zostałyby uznane za systemy wysokiego ryzyka, co nakładałoby na dostawców systemów AI dodatkowe obowiązki, a jednocześnie narzędzia służące do generowania materiałów *deepfake* pod ten rygor by nie podlegały. Uogólniając: mielibyśmy wówczas do czynienia z sytuacją, w której narzędzie służące do wykrywania przestępstw zostałoby uznane za narzędzie większego ryzyka niż narzędzia mogące posłużyć do ich popełnienia. Projekt zakładał w podpunkcie 2.3., że w przypadku systemów sztucznej inteligencji wysokiego ryzyka wymogi dotyczące danych wysokiej jakości, dokumentacji i identyfikowalności, przejrzystości, nadzoru ze strony człowieka, dokładności i solidności są absolutnie niezbędne, aby złagodzić stwarzane przez AI ryzyko dla praw podstawowych i bezpieczeństwa, którego nie obejmują istniejące ramy prawne⁴³. W ramach projektu w odniesieniu do niektórych szczególnych systemów sztucznej inteligencji, w tym zmanipulowanych cyfrowo obrazów lub filmów, zaproponowano wyłącznie minimalne obowiązki dotyczące przejrzystości. Zmiany, których w tym kontekście dokonano, należy ocenić pozytywnie, ponieważ wycofano się z pomysłu, aby narzędzie służące do wykrywania treści *deepfake* definiować jako narzędzie wysokiego ryzyka⁴⁴. Jednak warto zaznaczyć, że użyta w ramach art. 3 pkt 60 definicja pojęcia *deepfake* zawiera pewną lukę. W sytuacji gdy w ramach wygenerowanych lub zmanipulowanych przez sztuczną inteligencję obrazów, dźwięków lub materiałów filmowych użyto wizerunku osoby, definicja ograniczona jest wyłącznie do istniejących osób. Powoduje to sytuację, że gdyby

⁴¹ Załącznik III do wniosku dotyczącego rozporządzenia Parlamentu Europejskiego i Rady ustanawiającego zharmonizowane przepisy dotyczące sztucznej inteligencji (akt w sprawie sztucznej inteligencji) i zmieniającego niektóre akty ustawodawcze Unii (COM/2021/206 final).

⁴² Wniosek rozporządzenie Parlamentu Europejskiego i Rady ustanawiające zharmonizowane przepisy dotyczące sztucznej inteligencji (akt w sprawie sztucznej inteligencji) i zmieniające niektóre akty ustawodawcze Unii (COM/2021/206 final), zwany dalej Projektem.

⁴³ Punkt 2.3. uzasadnienia Projektu.

⁴⁴ Szczegółowy wykaz systemów AI wysokiego ryzyka, o których mowa w art. 6 ust. 2 zawiera załącznik III do rozporządzenia. Warunki, po których spełnieniu system zostanie uznany za system wysokiego ryzyka, opisuje ponadto art. 6 ust. 1 rozporządzenia.

w ramach wygenerowania przez sztuczną inteligencję – np. w nieprawdziwym materiale wideo – wykorzystano wizerunek osoby, która zmarła, to treść taka zgodnie z definicją zawartą w AI Act nie powinna zostać uznana za *deepfake*. Definicja zawiera słowa „treści (...), które przypominają istniejące osoby (...), które odbiorca mógłby niesłusznie uznać za autentyczne lub prawdziwe”⁴⁵. Pojęcie „istnieć” zgodnie ze Słownikiem języka polskiego oznacza „mieć miejsce w rzeczywistości”⁴⁶. Osoba, która odeszła już ze świata, bez wątpienia nie spełnia przesłanki istnienia. Jednocześnie w języku polskim używamy też pojęcia „przestać istnieć”, gdzie słowo „przestać” oznacza (w rozumieniu, w którym wypowiadamy się o jakimś stanie, procesie itp.): „skończyć się”⁴⁷. Zatem pojęcie „przestać istnieć” oznaczałoby: „skończyć proces, który ma miejsce w rzeczywistości”. Rozporządzenie 2024/1689 w tym samym miejscu w wersji anglojęzycznej używa słów „existing persons”, co wyklucza błąd tłumaczenia, a skłania do wniosku, iż doszło do niedopatrzenia w ramach tworzenia definicji⁴⁸.

Nie można uznać, że wykorzystanie wizerunku zmarłej osoby w ramach stworzenia nieprawdziwej treści wygenerowanej przez sztuczną inteligencję nie jest treścią typu *deepfake*. *De lege ferenda* definicja ta powinna zatem zostać uzupełniona i brzmieć w powyższym fragmencie: „treści (...), które przypominają istniejące obecnie lub w przeszłości osoby, (...), które odbiorca mógłby niesłusznie uznać za autentyczne lub prawdziwe”. Pozwoli to zlikwidować mogące w przyszłości wystąpić problemy interpretacyjne. Przykład ten nie jest rozumowaniem czysto akademickim, lecz sytuacją, jaka może zaistnieć w praktyce. Możliwe jest zdarzenie, w ramach którego – z wykorzystaniem autorytetu zmarłego przedsiębiorcy – wygenerowana zostanie treść, która będzie rzekomym ostatnim apelem, pomysłem na dalszy rozwój firmy lub upublicznieniem kompromitujących informacji. Należy pamiętać, że jedynym ograniczeniem w tym wypadku jest kreatywność twórcy takiej treści. Dodatkowo unormowania krajowe również są niekompletne, co powoduje, że problematyczne może być między innymi wskazanie osób legitymowanych czynnie do wystąpienia z powództwem o ochronę wizerunku osoby nieżyjącej, a pociągnięcie do odpowiedzialności osoby, która naruszyła prawo do wizerunku – niewykonalne⁴⁹.

⁴⁵ Art. 3 pkt 60 rozporządzenia.

⁴⁶ Słownik języka polskiego PWN, <https://sjp.pwn.pl/slowniki/istnie%C4%87.html> [dostęp: 20.12.2024].

⁴⁷ Słownik języka polskiego PWN, <https://sjp.pwn.pl/szukaj/przesta%C4%87.html> [dostęp: 20.12.2024].

⁴⁸ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act) (Text with EEA relevance), (OJ L, 2024/1689).

⁴⁹ O. Bodanka, *Naruszenie wizerunku przy wykorzystaniu technologii deepfake – analiza prawna i praktyczna*, „Opolskie Studia Administracyjno-Prawne” 2022, 20, 1, ss. 24 i 25.

W literaturze trafnie zauważa się, że w przypadku naruszenia wizerunku osoby nieżyjącej, szczególnie w sytuacji, gdy niemożliwe jest zidentyfikowanie sprawcy, osoby najbliższe mogą domagać się usunięcia tej treści od administratora portalu, na którym treść ta widnieje – zgodnie z art. 14 ustawy o świadczeniu usług drogą elektroniczną w zw. z art. 422 oraz art. 448 kodeksu cywilnego⁵⁰. Choć O. Bodanka słusznie wskazuje praktyczne trudności przy próbie ochrony wizerunku osoby zmarłej na podstawie uregulowań krajowych, to należy zaznaczyć, że problem jest znacznie szerszy, systemowy. Niewystarczające jest stwierdzenie, że ochrona wizerunku osób zmarłych jest ograniczona w ramach krajowego ustawodawstwa. Obecne unormowania prawne ujawniają przede wszystkim luki strukturalne na poziomie prawa unijnego (AI Act), ponieważ definicja „deepfake” zawarta w art. 3 pkt 60 rozporządzenia 2024/1689 dotyczy wyłącznie „istniejących osób” (ang. *existing persons*), co powoduje automatycznie wykluczenie sytuacji, w których wygenerowana z wykorzystaniem narzędzi AI treść odnosiłaby się bezpośrednio do osoby zmarłej. Opisany powyżej przypadek na szczeblu unijnym wpada więc niejako w próżnię prawną, ponieważ nie jesteśmy w stanie bez wątpliwości stwierdzić, czym – zgodnie z definicjami na szczeblu unijnym – tak wygenerowana treść jest. Rozważając rzecz dalej, należy stwierdzić, że tak ukształtowane rozumienie pojęcia „deepfake” powoduje wyłączenie przywołanych wyżej sytuacji z reżimu unijnych regulacji, w tym obowiązków dotyczących: przejrzystości, ujawnienia, że treść została wygenerowana za pomocą narzędzi AI, a także odpowiedzialności dostawcy systemu AI. Rozporządzenie 2024/1689 wprost zawiera informację, że obowiązek dotyczący przejrzystości, skutkujący koniecznością informowania o istnieniu w materiale tak wygenerowanych lub zmanipulowanych treści, dotyczy materiałów deepfake⁵¹. Dlatego obecną definicję należy poddać krytyce, następnie doprowadzając do jej poszerzenia, aby już na poziomie europejskim wykluczyć możliwe przyszłe wątpliwości interpretacyjne.

Warto powyższe rozważania zestawić z systemem prawnym obowiązującym na poziomie federalnym w Stanach Zjednoczonych, w którym Take It Down Act – zatwierdzony 19 maja 2025 roku – wprowadza pojęcie cyfrowego fałszerstwa (*digital forgery*)⁵². Dotyczy on jednak wyłącznie wizerunków wygenerowanych lub zmanipulowanych przez AI, które przedstawiają nagość lub treści o charakterze seksualnym, pomijając zagrożenia związane z oszustwami z wykorzystaniem technologii deepfake czy generowanymi w tym zakresie treściami dezinformującymi. Mimo że definicja zakresowo jest zdecydowanie węższa od zaproponowanej w art. 3 pkt 60 AI Act, to uni-

⁵⁰ Ibid., s. 25.

⁵¹ Motyw 134 preambuły rozporządzenia 2024/1689.

⁵² Tools to Address Known Exploitation by Immobilizing Technological Deepfakes on Websites and Networks Act, Statute at Large 139 Stat. 55 – Public Law No. 119-12 [dostęp: 19.05.2025].

ka sformułowania „existing persons” na rzecz przymiotnika „identifiable”. Obejmuje zatem swym zakresem każde intymne przedstawienie wizualne – niekoniecznie „istniejącej”, ale możliwej do zidentyfikowania osoby⁵³. Pozytywnie należy ocenić katalog kar, który w przypadku publikacji materiału typu *deepfake* dotyczącej osoby dorosłej przewiduje karę grzywny, karę pozbawienia wolności do lat 2 lub obie te kary łącznie⁵⁴. W przypadku publikacji dotyczącej osoby małoletniej górna granica pozbawienia wolności rośnie do lat 3⁵⁵. Przepisy te w istocie obejmują powszechny problem, który należałoby uregulować również w ramach polskiego ustawodawstwa, ale jednocześnie zauważalne jest tu traktowanie zagrożeń gospodarczych jako społecznie mniej istotnych od intymnych treści wygenerowanych z użyciem technologii *deepfake*.

Również Republika Korei w swym systemie prawnym zawiera obostrzenia dotyczące generowania treści z wykorzystaniem technologii *deepfake*⁵⁶. Art. 14-2 *expressis verbis* kryminalizuje edycję, łączenie lub przetwarzanie zdjęć, nagrań wideo lub audio. Ponadto penalizacji poddano rozpowszechnianie tych materiałów albo ich kopii. Wątpliwości budzi to, czy karze podlega rozpowszechnianie takiego materiału po śmierci tej osoby wbrew jej woli⁵⁷. Znamienna jest również surowość kar, gdzie górną granicę wyznaczono na 7 lat prac przymusowych połączonych z pozbawieniem wolności albo karę grzywny nieprzekraczającą 50 milionów wonów południowokoreańskich. Na tle regulacji unijnych – które w ramach art. 50 ust. 4 wprowadzają się do iluzorycznego obowiązku ujawniania przez sprawcę, iż treść została wygenerowana z użyciem sztucznej inteligencji – rozwiązanie koreańskie wyróżnia się penalizacyjną precyzją. Podkreślenia wymaga fakt, że regulacje koreańskie – podobnie jak *Take It Down Act* – skoncentrowane są na ochronie ofiar przestępstw seksualnych, z pominięciem elementu zagrożeń gospodarczych.

⁵³ *Take It Down Act*, Public Law 119-12, § 2, nowelizujący sekcję 223 *Communications Act of 1934*, 47 U.S.C. § 223(h)(1)(B).

⁵⁴ *Take It Down Act*, Public Law 119-12, § 2, nowelizujący sekcję 223 *Communications Act of 1934*, 47 U.S.C. § 223(h)(4).

⁵⁵ *Ibid.*

⁵⁶ *Act on Special Cases Concerning the Punishment of Sexual Crimes*, Act No. 20459, amended Oct. 16, 2024, art. 14-2, Republic of Korea, Korea Legislation Research Institute, https://elaw.klri.re.kr/eng_service/lawView.do?hseq=68812&lang=ENG [dostęp: 10.05.2026].

⁵⁷ Anglojęzyczne tłumaczenie wprost zawiera wyrażenie „after the death”. Nasuwa się pytanie, czy oficjalne tłumaczenie mogłoby pomylić sformułowanie „after the death” z pojęciem „subsequently”. Wątpliwość wynika z faktu, że koreański termin 사후에 (użyty w art. 14-2 ust. 2) rodzi wątpliwości interpretacyjne i może być rozumiany jako: po fakcie, następnie oraz po śmierci; wobec braku zapisu ‘hanja’ w tekście ustawy oraz rozstrzygających wyjaśnień, precyzyjne ustalenie znaczenia wymagałoby poznania oficjalnego stanowiska ustawodawcy koreańskiego.

Wyzwania dotyczące odpowiedzialności cywilnej za czyny dokonane z wykorzystaniem systemów AI

Ostatnim, lecz nie mniej ważnym aspektem jest egzekwowanie przyjętych przepisów. Rozdział XII rozporządzenia 2024/1689 przewiduje katalog kar pieniężnych w przypadku nierespektowania przepisów wynikających z AI Act. Najsurowszymi karami obłożono działania, które stanowiłyby naruszenie art. 5 rozporządzenia 2024/1689, wskazującego zakazane praktyki w zakresie AI. Nieprzestrzeganie tego przepisu podlega administracyjnej karze pieniężnej w wysokości do 35 000 000 EUR lub – jeżeli sprawcą naruszenia jest przedsiębiorstwo – w wysokości do 7 % jego całkowitego rocznego światowego obrotu z poprzedniego roku obrotowego, w zależności od tego, która z tych kwot jest wyższa⁵⁸. Ponadto państwa członkowskie mają obowiązek ustanowienia przepisów dotyczących kar i innych środków egzekwowania prawa, które mogą również obejmować ostrzeżenia i środki niepieniężne, mających zastosowanie w przypadku naruszeń przepisów rozporządzenia 2024/1689 przez operatorów; mają też obowiązek podejmowania wszelkich niezbędnych środków w celu zapewnienia ich właściwego i skutecznego wykonywania, z uwzględnieniem wytycznych wydanych przez Komisję, zgodnie z jego art. 96⁵⁹.

Brak jednolitego unijnego reżimu odpowiedzialności cywilnej za szkody wyrządzone z użyciem systemów AI wymusza konieczność oparcia się na przepisach krajowych⁶⁰. Równocześnie krajowy porządek prawny nie przewiduje *explicite* cywilnej odpowiedzialności za czyny *ex delicto* popełnione z użyciem systemów sztucznej inteligencji. Na poziomie unijnym uzupełnieniem jest dyrektywa Parlamentu Europejskiego i Rady (UE) 2024/2853 z 23 października 2024 r. w sprawie odpowiedzialności za produkty wadliwe⁶¹. W ramach art. 4 dyrektywy 2024/2853 włączono do definicji produktu oprogramowanie⁶². Jednocześnie zakłada ona, aby odpowiedzialnym podmiotem gospodarczym był producent wadliwego produktu, wadliwej części składowej lub – w przypadku producenta produktu lub jego części składowej, który ma siedzibę poza UE oraz bez uszczerbku dla odpowiedzialności tego producenta – odpowiedzialność ponosił importer wadliwego produktu lub części składo-

⁵⁸ Wymiar poszczególnych kar opisuje art. 99 ust. 3, 4 oraz 5 rozporządzenia 2024/1689.

⁵⁹ Art. 99 ust. 1 rozporządzenia 2024/1689.

⁶⁰ Projekt dyrektywy w sprawie odpowiedzialności za AI (AI Liability Directive, COM/2022/496 final), który miał uzupełnić AI Act, został wycofany przez Komisję Europejską w 2025 roku wobec braku perspektyw na osiągnięcie porozumienia legislacyjnego.

⁶¹ Dyrektywa Parlamentu Europejskiego i Rady (UE) 2024/2853 z 23 października 2024 r. w sprawie odpowiedzialności za produkty wadliwe i uchylenia dyrektywy Rady 85/374/EWG (Dz.U. UE L 2024.2853 z późn. zm.), dalej: dyrektywa 2024/2853.

⁶² Dyrektywa 2024/2853 definiuje produkt jako rzecz ruchomą, nawet jeżeli jest ona zintegrowana lub wzajemnie połączona z inną rzeczą ruchomą lub nieruchomą; obejmuje to energię elektryczną, cyfrowe pliki produkcyjne, surowce i oprogramowanie.

wej, przedstawiciel, który został przez producenta upoważniony lub dostawca usług realizacji zamówień, gdy brak jest importera mającego siedzibę w Unii oraz upoważnionego przedstawiciela⁶³. Termin transpozycji omawianej dyrektywy przypada na 9 grudnia 2026 r., co oznacza, że do tego czasu krajowy system prawny jest pozbawiony wyraźnej podstawy do stosowania odpowiedzialności za produkt wadliwy wobec wykorzystania systemów AI służących generowaniu treści *deepfake*⁶⁴.

Do 30 czerwca 2026 roku polski ustawodawca nie dokonał implementacji dyrektywy 2024/2853 – czego skutkiem jest obecnie obowiązujący art. 449¹ k.c., który w § 2 obejmuje definicję produktu, natomiast w § 1 zaznacza odpowiedzialność producenta. Jednocześnie przepis wskazuje tylko taki podmiot, który wytworzył produkt w ramach prowadzonej działalności gospodarczej⁶⁵. Obecnie obowiązująca norma prawna stanowi implementację dyrektywy 85/374/EWG, bezpośrednio poprzedzającej dyrektywę 2024/2853⁶⁶. Nawet uwzględniając zaktualizowaną już definicję pojęcia „produkt”, obejmującą oprogramowanie, wątpliwości budzi zakres odpowiedzialności producenta w zakresie wygenerowania za pomocą systemu AI materiału *deepfake*, który miał zamiar lub faktycznie wyrządził szkodę. Zasadniczo udostępnienie systemu nie powinno przesądzać o odpowiedzialności producenta, a wymagana winna być weryfikacja, czy producent dochował należytej staranności przy projektowaniu wymaganych przepisami prawa mechanizmów zabezpieczających przed niewłaściwym użyciem systemu przez użytkownika końcowego.

W ramach wyzwań związanych z zaistnieniem możliwych szkód dla przedsiębiorców – ale również osób fizycznych nieprowadzących działalności gospodarczej – związanych z nieetycznym użyciem narzędzi AI zauważono, że *de lege lata* AI Act przewiduje mechanizm, zgodnie z którym podmiot chcący wygenerować materiał *deepfake* powinien w sposób wyraźny i jasny ujawnić, że przedstawione treści zostały sztucznie wygenerowane lub zmanipulowane⁶⁷. Rozwiązanie to opiera się na naiwnym założeniu, że podmiot działający *animo nocendi* dobrowolnie ujawni poszkodowanemu fakt popełnienia czynu niedozwolonego. Takie założenie jest sprzeczne z logiką, gdyż nikt, kto działa w złej wierze, nie ułatwia świadomie ustalenia istoty swojego czynu⁶⁸.

⁶³ Art. 8 ust. 1 dyrektywy 2024/2853.

⁶⁴ Zakres oraz termin transpozycji dyrektywy 2024/2853 opisuje jej art. 22.

⁶⁵ G. Karaszewski, komentarz do art. 449¹ [w:] Kodeks cywilny. Komentarz aktualizowany, red. P. Nazaruk, Gdańsk 2024.

⁶⁶ Dyrektywa Rady z 25 lipca 1985 r. w sprawie zbliżenia przepisów ustawowych, wykonawczych i administracyjnych państw członkowskich dotyczących odpowiedzialności za produkty wadliwe (Dz.U. UE L 1985.210.29 z późn. zm.).

⁶⁷ Zob. motyw 134 Preambuły AI Act oraz art. 50 ust. 4 AI Act.

⁶⁸ P. Michalik, P. Zaremba, *Deepfake jako dowód w procesie cywilnym – rozważania na tle polskiego postępowania cywilnego*, „Prawo i Więź” 2025, 57, 4, s. 639.

Niewątpliwie istotne jest, aby obowiązujące przepisy były respektowane, a ich wykonywanie skutecznie kontrolowane i egzekwowane. Jednakże należy pamiętać, że nadmierna regulacja i zbytne obciążenie przedsiębiorców obowiązkami prawnymi negatywnie wpływa na innowacyjność firm, która jest kluczowa dla rozwoju gospodarki – w szczególności w sektorze nowych technologii. Według raportu sporządzonego w grudniu 2022 roku pod tytułem AI Act Impact Survey wynika, że dwie trzecie startupów spodziewa się negatywnego wpływu rozporządzenia 2024/1689 na innowacje w zakresie sztucznej inteligencji w Europie⁶⁹. *De lege ferenda* należałoby więc maksymalnie uprościć kwestie dokumentacyjne dla przedsiębiorców. Obecna sytuacja – w której załącznik VIII opisuje informacje przekazywane przy rejestracji systemów AI wysokiego ryzyka zgodnie z art. 49 rozporządzenia 2024/1689, gdzie załącznik ten jest dodatkowo podzielony na trzy odrębne sekcje, następnie załącznik IX, który opisuje informacje przedkładane przy rejestracji systemów AI wysokiego ryzyka wymienionych w załączniku III do rozporządzenia w odniesieniu do testów w warunkach rzeczywistych zgodnie z art. 60 AI Act – nie ułatwia rozwoju przedsiębiorstw zajmujących się sztuczną inteligencją⁷⁰. Należy pamiętać, że zagrożenia gospodarcze wynikające z rozwoju nowoczesnych technologii mogą dotyczyć przedsiębiorców w ogólności – przez używanie sztucznej inteligencji w sposób nieetyczny, w celu osiągnięcia konkretnego celu dezinformującego, dyskredytującego-ośmieszającego czy też cyberprzestępczego – ale też mogą dotyczyć wyłącznie przedsiębiorstw zajmujących się rozwojem nowych technologii, dla których nadmierna biurokracja może stać się impulsem do przeniesienia działań biznesowych na inną część globu, w której regulacje prawne będą zdecydowanie mniej uciążliwe.

Niemniej jednak istnieją pewne zagrożenia skłaniające do dalszego regulowania podatnych na nie materii. Odnosi się to przede wszystkim do przepisów krajowych, gdzie należałoby wprowadzić odrębne przepisy w ramach prawa karnego, które uzupełniłyby lukę wynikającą z burzliwego i wciąż przyspieszającego rozwoju technologicznego. Zagrożenia te dotyczą przy tym wykorzystania technologii *deepfake* zarówno w sferze gospodarczej, jak w kontekście naruszenia dóbr konkretnych jednostek, takich jak godność, cześć, dobre imię. Jak słusznie zauważa A. Ziobroń: niewystarczający zakres regulacji potencjalnie szkodliwego społecznie zjawiska może doprowadzić do osłabienia funkcji sprawiedliwościowej prawa karnego⁷¹. Dlatego prace nad nowelizacją kodeksu karnego w tym zakresie powinny nastąpić tak szybko, jak to tylko możliwe.

⁶⁹ AI Act Impact Survey, https://www.appliedai.de/uploads/files/AI-Act-Impact-Survey_Slides_Dec12.2022.pdf [dostęp: 27.12.2024].

⁷⁰ Załącznik VIII oraz IX do rozporządzenia.

⁷¹ A. Ziobroń, *Deepfake a prawo...*, s. 235.

Podsumowanie

W ramach niniejszej analizy wskazano zagrożenia gospodarcze wynikające z wykorzystania technologii *deepfake* oraz ich implikacje prawne i społeczne. W kontekście gospodarczym zidentyfikowano kluczowe zagrożenia, takie jak:

- a) dezinformacja – rozpowszechnianie fałszywych informacji w celu manipulacji rynkami kapitałowymi lub szkodenia reputacji firm;
- b) naruszenie dóbr osobistych – ataki *deepfake* skierowane na liderów biznesu, co może prowadzić do utraty zaufania ze strony społeczeństwa, a w konsekwencji do strat finansowych przedsiębiorstw;
- c) cyberoszustwa – wyłudzenie środków z wykorzystaniem spreparowanych treści, np. podszywanie się pod znanych biznesmenów w celu zachęcenia do fałszywych inwestycji;
- d) inne cyberprzestępstwa – wykorzystanie *deepfake*’ów do zwiększania skuteczności ataków phishingowych, generowania fałszywych recenzji lub innego manipulowania opinią publiczną.

Odnosząc się do głównego problemu badawczego przedstawionego w niniejszym artykule, należy stwierdzić, że obecne mechanizmy prawne w UE – ze szczególnym uwzględnieniem AI Act – stanowią duży krok naprzód w sferze uregulowań związanych z technologicznym rozwojem oraz wynikających z tego zagrożeń. Należy jednak zaznaczyć, że z przeprowadzonej analizy jednoznacznie wynika, iż regulacje te nie są wystarczające, by w sposób właściwy chronić przed gospodarczymi skutkami nadużyć dokonanych z wykorzystaniem narzędzi AI. Rozporządzenie 2024/1689 w obecnym brzmieniu nie przewiduje stworzenia stosownych mechanizmów pozwalających na egzekwowanie odpowiedzialności za treści generowane przez systemy sztucznej inteligencji, które nie mają osobowości prawnej, a jednocześnie skutki ich wykorzystywania są realne. Słusznym rozwiązaniem było wprowadzenie definicji legalnej pojęcia „*deepfake*”, jednak definicja ta ma lukę, ponieważ nie obejmuje sytuacji, gdy wykorzystano wizerunek osoby zmarłej – np. niedawno zmarłego prezesa zarządu. Ograniczenie definicji wyłącznie do istniejących osób może zostać świadomie wykorzystane przez sprawców chcących posłużyć się wizerunkiem osoby zmarłej. Zidentyfikowana luka nie powinna jednak być uzupełniana przez wykładnię, lecz wymaga interwencji legislacyjnej – ze względu na jej strukturalny, a nie redakcyjny charakter. Dodatkowo mechanizm wynikający z art. 50 ust. 4 AI Act, zakładający dobrowolne informowanie przez sprawcę oszustwa, że korzystał w jego trakcie z systemów AI, jest sprzeczny z elementarną logiką. Przepis ten jedynie pozoruje realną ochronę, zamiast rzeczywiście ją tworzyć. Zbytняя ogólność przepisów dotyczących ochrony przed nieetycznie wykorzystanym narzędziem AI, pozwalająca tworzyć oszukańcze bądź dezinformacyjne treści *deepfake*, ogranicza skuteczność uregulowań prawnych wynikających z AI Act.

Z drugiej strony nadmiar biurokratycznych i innych obowiązków nałożony na przedsiębiorców może zniechęcać ich do podejmowania technologicznie innowacyjnych działań – co bezpośrednio wpłynie będzie na konkurencyjność europejskiego rynku w porównaniu z resztą świata. Dlatego też wnioskiem płynącym z badań jest potrzeba uzupełnienia przepisów unijnych o rozszerzenie definicji pojęcia *deepfake*; zarazem jednak konieczne staje się ograniczenie i uproszczenie wymagań administracyjnych dla przedsiębiorców działających w sektorze sztucznej inteligencji. Dopiero synteza precyzyjnie ukształtowanych przepisów z jednoczesnym odbiurokratyzowaniem wymagań narzuconych na europejskie podmioty gospodarcze przyniesie odpowiednie rozwiązanie – łączące przeciwdziałanie zagrożeniom gospodarczym z zachowaniem odpowiedniego poziomu innowacji i rozwoju. Jednoznacznym wnioskiem z przeprowadzonej analizy jest potrzeba pilnej rewizji obowiązujących ram regulacyjnych, począwszy od definicji normatywnych, a skończywszy na mechanizmach egzekwowania odpowiedzialności prawnej podmiotów używających systemów AI w sposób naruszający obowiązujące przepisy. Jednocześnie jednym z większych wyzwań na szczelnie unijnym pozostaje ukształtowanie przepisów w taki sposób, aby zneutralizować jak najwięcej zagrożeń gospodarczych związanych z rozwojem sztucznej inteligencji, ale jednocześnie zrealizować ten cel tak, by inne państwa – w szczególności Stany Zjednoczone Ameryki czy Chińska Republika Ludowa – nie zdominowały kosztami Europy światowego rozwoju technologicznego. Rozszerzenie definicji *deepfake* o treści wygenerowane z wykorzystaniem wizerunków osób zmarłych, zastąpienie art. 50 ust. 4 skuteczną sankcją cywilną wraz z legitymacją czynną poszkodowanego wobec sprawcy wywołującego szkodę oraz konieczność implementacji dyrektywy 2024/2853 nie są bynajmniej postulatami nadregulacyjnymi, lecz koniecznym uzupełnieniem obecnie obowiązujących przepisów w celu zwiększenia ich skuteczności i precyzji.

Abstrakt

Przedmiotem artykułu jest analiza wpływu sztucznej inteligencji, a w szczególności technologii *deepfake* na gospodarkę. Praca ma za cel zweryfikowanie, czy obecne unormowania prawne na szczelnie europejskim w wystarczającym stopniu chronią przed zagrożeniami, jakie stawia przed nami rozwój technologiczny. Przedstawiono przypadki, w których można wykorzystać technologię *deepfake* w celu wyłudzenia środków finansowych czy podjęcia innych nieetycznych działań. Kluczowe zagadnienia obejmują rolę unijnych regulacji – takich jak akt w sprawie sztucznej inteligencji – w celu przeciwdziałania zagrożeniom gospodarczym. Ponadto zwrócono uwagę na konieczność zmian legislacyjnych, które uzupełnią występujące obecnie luki prawne, a jednocześnie zmniejszą

szą zakres obowiązków narzucanych na przedsiębiorców. Wskazano, że unormowania prawne dotyczące sztucznej inteligencji są niezbędne – z jednoczesnym zauważeniem, że ogromnym wyzwaniem jest utrzymanie odpowiedniego poziomu innowacyjności przedsiębiorstw w Europie.

Słowa kluczowe: deepfake, gospodarka, akt w sprawie sztucznej inteligencji, zagrożenia, sztuczna inteligencja (AI).

BIBLIOGRAFIA

Bodanka O., Naruszenie wizerunku przy wykorzystaniu technologii deepfake – analiza prawna i praktyczna, „Opolskie Studia Administracyjno-Prawne” 2022, 20, 1.

Bukowski M., Zwalczanie cyberprzestępczości ekonomicznej przy wykorzystaniu sztucznej inteligencji (AI), „Przegląd Policyjny” 2023, 151, 3.

Dąbrowska I., Deepfake: nowy wymiar internetowej manipulacji, „Zarządzanie Mediami” 2020, 8, 2.

Kiełpiński K., Deepfake jako narzędzie do przekazywania informacji fałszywej i domniemanej. Analiza prawnokarna i cybernetyczna, „Kwartalnik Krajowej Szkoły Sądownictwa i Prokuratury” 2023, 3, 51.

Kira B., When Non-Consensual Intimate Deepfakes Go Viral: The Insufficiency of the UK Online Safety Act, ‘Computer Law & Security Review’ 2024, 54.

Kodeks cywilny. Komentarz aktualizowany, P. Nazaruk (red.), Gdańsk 2024.

Kumari D., Deepfake Technology and Legal Issues, ‘LawFoyer International Journal of Doctrinal Legal Research’ 2024, 2, 1.

Lai L., Świerczyński M. (red.), Prawo sztucznej inteligencji, Warszawa 2020.

Michalik P., Zaremba P., Deepfake jako dowód w procesie cywilnym – rozważania na tle polskiego postępowania cywilnego, „Prawo i Więź” 2025, 57, 4.

Mierzejewski D., Nawrotkiewicz J., W obliczu wspólnej polityki UE wobec Chin: „mission impossible”?, „Pulaski Policy Papers” 2024, XI.

Piecuch A., Sztuczna inteligencja w perspektywie społecznej, „Edukacja Ustawiczna Dorosłych” 2023, 123, 4.

Skutki i zagrożenia cywilizacji informacyjnej, K.A. Nawrot, K. Prandecki (red.), Komitet Prognoz PAN 2023.

Spałek D., Rzymowski J., Jak nie stosować AI Act na podstawie definicji systemu sztucznej inteligencji, „Prawo Nowych Technologii” 2024, 2.

Stewart I., *Significant Figures Lives and Works of Trailblazing Mathematics*, New York 2017 [wyd. pol. *Krótką historia wielkich umysłów. Genialni matematycy i ich arcydzieła*, tłum. U. Seweryńska i M. Seweryński, Warszawa 2019].

Turing A., *Computing Machinery and Intelligence*, 'Mind' 1950, 59, 236.

Vera G.G., Deep fake (!) – postęp technologiczny a prawo karne, „Zeszyty Naukowe Uniwersytetu Rzeszowskiego – Seria Prawnicza” 2024, 1, 44.

Widło J., Odpowiedzialność za szkodę wyrządzoną przez pojazdy autonomiczne w świetle zmian prawa UE, „Forum Prawnicze” 2025, 2.

Zielińska A., Korzystanie z treści chronionych a sztuczna inteligencja – jak daleko sięga odpowiedzialność dostawcy usług udostępniania treści online w świetle art. 17 dyrektywy 2019/790?, „Studia de Cultura” 2022, 14, 2.

Ziobroń A., Deepfake a prawo karne. Uwagi „de lege lata” i „de lege ferenda” dotyczące fałszywej pornografii, „Studenckie Prace Prawnicze, Administratywistyczne i Ekonomiczne” 2021, 37.

Źródła internetowe

AI Act Impact Survey, https://www.appliedai.de/uploads/files/AI-Act-Impact-Survey_Slides_Dec12.2022.pdf [dostęp: 27.12.2024].

Brandon J., Terrifying high-tech porn: Creepy 'deepfake' videos are on the rise, <https://www.foxnews.com/tech/terrifying-high-tech-porn-creepy-deepfake-videos-are-on-the-rise> [dostęp: 7.12.2024].

Crime as a service (CaaS), <https://www.comcert.pl/slownik/crime-as-a-service-caas/> [dostęp: 17.12.2024].

FBI, Internet Crime Report 2022, https://www.ic3.gov/AnnualReport/Reports/2022_ic3report.pdf [dostęp: 24.10.2025].

FBI, Internet Crime Report 2024, https://www.ic3.gov/AnnualReport/Reports/2024_IC3Report.pdf [dostęp: 24.10.2025].

FBI, Internet Crime Report 2025, https://www.ic3.gov/AnnualReport/Reports/2025_IC3Report.pdf [dostęp: 8.05.2026].

Jakubczak D., Deepfake. Oszustwo na sklonowanego dyrektora, <https://www.dw.com/pl/deepfake-oszustwo-na-sklonowanego-dyrektora/a-69116765> [dostęp: 16.12.2024].

Lista osób, których wizerunek przestępcy wykorzystali w oszustwach deepfake, <https://nask.pl/aktualnosci/nask-ostrega-przestepcy-tworza-coraz-bardziej-pomyslowne-oszustwa-deepfake/> [dostęp: 12.12.2024].

NASK ostrzega: przestępcy tworzą coraz bardziej pomysłowe oszustwa deepfake, <https://nask.pl/aktualnosci/nask-ostrega-przestepcy-tworza-coraz-bardziej-pomyslowne-oszustwa-deepfake/> [dostęp: 12.12.2024].

Online Safet Act 2023, c. 50, Zjednoczone Królestwo, <https://www.legislation.gov.uk/ukpga/2023/50/contents/enacted> [dostęp: 10.05.2026].

Słownik języka polskiego PWN, <https://sjp.pwn.pl/slowniki/istnie%C4%87.html> [dostęp: 20.12.2024].

Słownik języka polskiego PWN, <https://sjp.pwn.pl/szukaj/przesta%C4%87.html> [dostęp: 20.12.2024].

Unni A., What You Need To Know About Crime As A Service (CSaaS), <https://www.stickmancyber.com/cybersecurity-blog/what-you-need-to-know-about-crime-as-a-service-csaas> [dostęp: 17.12.2024].