

Jakub Jurowicz

Uniwersytet Jagielloński
Wydział Prawa i Administracji
ORCID: [0009-0002-7113-5992](https://orcid.org/0009-0002-7113-5992)

Małgorzata Paja

Uniwersytet Jagielloński
Wydział Prawa i Administracji
ORCID: [0009-0003-0779-9636](https://orcid.org/0009-0003-0779-9636)

Zwielokrotnianie utworów w procesach tworzenia i funkcjonowania dużych modeli językowych

Streszczenie

W odniesieniu do tzw. dużych modeli językowych (ang. *large language models*, LLM) jednym z głównych zagadnień jest zagrożenie dla wyłącznych praw twórców, a w szczególności – praw autorskich. Wspomniane zagrożenie materializuje się zarówno na etapie powstawania LLM, jak i późniejszego funkcjonowania takiego modelu. W przedmiocie korzystania z cudzej twórczości spory w doktrynie prawa toczą się przede wszystkim w odniesieniu do zakresu dozwolonego użytku utworów w ramach eksploracji tekstów i danych, w odniesieniu do tzw. zjawiska zapamiętywania danych treningowych przez LLM oraz wobec zjawiska tzw. regurgitacji, tj. „zwracania” fragmentów chronionych danych treningowych przez LLM w generowanych treściach. Wskazane spory podsycało procedowanie najnowszej legislacji dotyczącej sztucznej inteligencji, tj. przede wszystkim Aktu o sztucznej inteligencji, którego kolejne wersje pojawiały się w I kwartale 2024 r. i którego ostateczna wersja została uchwalona 13 czerwca 2024 r. – uwzględniając odesłania do przepisów o eksploracji tekstów i danych na gruncie Dyrektywy o prawie autorskim na jednolitym rynku cyfrowym. Celem niniejszego tekstu jest analiza poszczególnych etapów powstawania i funkcjonowania LLM oraz weryfikacja zachodzącego w ich trakcie zwielokrotniania chronionej twórczości w kontekście ewentualnego stosowania przepisów o dozwolonym użytku.

Słowa kluczowe

AI, utwór, eksploatacja, dozwolony użytek, eksploracja, tekstów, danych, TDM

Wstęp

Dynamiczny rozwój sztucznej inteligencji (dalej: „AI”) spowodował równie gwałtowny wzrost liczby problemów natury prawnej. W odniesieniu do modeli AI

ogólnego przeznaczenia¹ jednym z głównych zagadnień jest zagrożenie dla wyłącznych praw twórców, a w szczególności majątkowych praw autorskich. Wspomniane zagrożenie materializuje się co do zasady na dwóch płaszczyznach, tj. na etapie trenowania modelu AI oraz na etapie tworzenia generacji przez AI. Autorzy, dla zwiększenia przejrzystości dalszej analizy, wskazują w tym miejscu, że na potrzeby niniejszego artykułu, mając na myśli modele AI ogólnego przeznaczenia, będą posługiwać się terminem „LLM” (ang. *large language models*)² – jako że rozważania poczynione w takcie analizy odnoszą się, jak wskazuje tytuł, właśnie do dużych modeli językowych.

Powyższe zagadnienie wróciło na agendę doktrynalnej dyskusji z nową mocą, gdy wraz z kolejnymi uchwalonymi wersjami Aktu o sztucznej inteligencji 13 marca 2024 r. Parlament Europejski przyjął swoją wersję tekstu Aktu o sztucznej inteligencji³, następnie 14 maja 2024 r. swoją wersję uchwaliła Rada Europejska⁴, a końcowe brzmienie AI Act ukształtowano w wersji z 13 czerwca 2024 r. W toku prac nad kolejnymi wersjami AI Act prawodawca europejski zdecydował się zaadresować problemy dotyczące trenowania LLM na dziełach objętych cudzymi prawami autorskimi (por. ostateczna treść art. 53 ust. 1 lit. c) AI Act), ujmując pośród obowiązków dostawców takich modeli obowiązek wprowadzenia polityki zapewnienia zgodności

¹ Model AI, w tym model AI trenowany dużą ilością danych z wykorzystaniem nadzoru własnego na dużą skalę, który wykazuje znaczną ogólność i jest w stanie kompetentnie wykonywać szeroki zakres różnych zadań, niezależnie od sposobu, w jaki model ten jest wprowadzany do obrotu, i który można zintegrować z różnymi systemami lub aplikacjami niższego szczebla – z wyłączeniem modeli AI, które są wykorzystywane na potrzeby działań w zakresie badań, rozwoju i tworzenia prototypów przed wprowadzeniem ich do obrotu – za definicją z art. 3 pkt 63 rozporządzenia Parlamentu Europejskiego i Rady (UE) 2024/1689 z dnia 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, z 12.7.2024) (dalej: „AI Act”).

² I. Emanuilov, T. Margoni, *Memorisation in generative models and EU copyright law: an interdisciplinary view*, Kluwer Copyright Blog, 26.03.2024, <https://copyrightblog.kluweriplaw.com/2024/03/26/memorisation-in-generative-models-and-eu-copyright-law-an-interdisciplinary-view/> [dostęp 21.05.2024].

³ Rezolucja ustawodawcza Parlamentu Europejskiego z dnia 13 marca 2024 r. w sprawie 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)0206 – C9-0146/2021 – 2021/0106(COD)), https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_PL.pdf [dostęp 17.05.2024].

⁴ Propozycja Rozporządzenia Parlamentu Europejskiego i Rady 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) w wersji przyjętej 14 maja 2024 r., 2021/0106(COD), PE-CONS 24/24, <https://data.consilium.europa.eu/doc/document/PE-24-2024-INIT/pl/pdf> (dalej: „AI Act”) [dostęp 21.05.2024].

z unijnym prawem autorskim, w szczególności z myślą o identyfikacji i przestrzeganiu, w tym poprzez najnowocześniejsze technologie, zastrzeżenia praw wyrażonego zgodnie z art. 4 ust. 3 dyrektywy (UE) 2019/790⁵.

W toku tworzenia LLM do zwielokrotniania utworów chronionych prawem autorskim może dochodzić na kilku etapach:

1) na etapie tworzenia LLM, tj.:

- a) w trakcie pozyskiwania danych zawierających utwory – aby wytrenowanie LLM było możliwe, konieczne jest zgromadzenie znacznej ilości danych. Odbywa się to na wiele sposobów – np. poprzez wewnętrzne gromadzenie danych (wydobywanie ich z zasobów twórcy modelu), *crowdsourcing* (pozyskiwania danych poprzez przekazywanie ich przez inne osoby) lub poprzez zautomatyzowane gromadzenie danych (np. w drodze tzw. *web scrapingu*, polegającego na wydobywaniu danych ze stron internetowych za pomocą różnych narzędzi informatycznych)⁶;
- b) podczas przetwarzania takich utworów w celu uzyskania odpowiedniej jakości danych treningowych dla trenowanego LLM – czynności te obejmują sprawdzanie danych zgromadzonych do treningu pod kątem adekwatności i kompletności, a następnie oczyszczanie danych i formatowanie ich na potrzeby szkolenia⁷;

2) w ramach procesu zwanego zapamiętywaniem lub memoryzacją (ang. *memorisation*) odbywającego się w niektórych LLM – z prowadzonych badań⁸ wynika, że niektóre LLM budowane w oparciu o mechanizmy *deep learning* są podatne na zapamiętywanie (tj. przechowywanie) fragmentów danych treningowych. Przyjmuje się, że tak jak ludzki mózg zapamiętuje fragmenty informacji, aby się uczyć, tak samo robią też LLM;

3) w wyniku generowania treści przez AI po wprowadzeniu przez użytkownika promptu lub promptów – w procesie korzystania z LLM możliwa jest

⁵ Dyrektywa Parlamentu Europejskiego i Rady (UE) 2019/790 z dnia 17 kwietnia 2019 r. w sprawie prawa autorskiego i praw pokrewnych na jednolitym rynku cyfrowym oraz zmiany dyrektyw 96/9/WE i 2001/29/WE, Dz. Urz. UE L 130 z 17.05.2019 r., s. 92–125 (dalej: „dyrektywa DSM”).

⁶ J. Tan, *What Is AI Model Training?*, 30.04.2024, <https://engage-ai.co/what-ai-model-training/> [dostęp 21.05.2024].

⁷ D. O’Keefe, *How Does AI Model Training Work?*, apian, 24.01.2014, <https://appian.com/blog/acp/ai/how-does-ai-model-training-work.html> [dostęp 21.05.2024]; AILabs, *Data Cleaning and AI Model Training in Algorithmic Training*, <https://www.ailabs.global/blog/data-cleaning-and-ai-model-training-in-algorithmic-training> [dostęp 21.05.2024].

⁸ N. Carlini *et al.*, *Extracting Training Data from Large Language Models*, <https://www.usenix.org/conference/usenixsecurity21/presentation/carlini-extracting> [dostęp 21.05.2024].

sytuacja, w której model „zwraca” w swoich generacjach (tj. w outputcie) mniejsze lub większe fragmenty danych treningowych – a więc także utworów (np. tekstów lub grafik)⁹.

Celem niniejszego tekstu jest omówienie zagadnienia zwielokrotniania utworów w kontekście tworzenia i funkcjonowania LLM oraz próba rozstrzygnięcia, czy *de lege lata* takie korzystanie z cudzych utworów wymaga zgody podmiotu uprawnionego, czy też nie wymaga takiej zgody.

1. Kwestie techniczne i terminologiczne

1.1. Trenowanie LLM

Przed przystąpieniem do szczegółowych rozważań w przedmiocie zwielokrotniania utworów (1) w procesie tworzenia LLM, (2) dokonującego się w LLM poprzez ich zapamiętywanie oraz (3) zwielokrotniania odbywającego się w trakcie generowania treści przez AI, a będącego wynikiem zapamiętywania chronionych i w konsekwencji – tzw. regurgitacji treści wchodzącej w skład danych treningowych LLM¹⁰, należy wyjaśnić zagadnienia wstępne, których zrozumienie jest niezbędne dla poprawnego śledzenia toku dalszego wywodu.

LLM potrafią generować np. teksty, obrazy i inne treści. Opracowanie i trenowanie takich modeli wymagają jednak dostępu do ogromnych ilości danych. Dane te są przez algorytmy oparte na uczeniu maszynowym (a konkretniej – jego podkategorii, tj. uczeniu głębokim, ang. *deep learning*) analizowane w procesie zwanym trenowaniem lub uczeniem. Wspomniane algorytmy robią to konsekwentnie poprzez uogólnianie dostarczonych danych i wychwytywanie występujących w dostarczonych danych wzorców i relacji¹¹. Dzięki temu trenowany model może następnie stosować wyuczone na wprowadzonych danych wzorce i zależności, analizując nowe dane, które wcześniej nie były modelowi znane¹² (np. dane wejściowe wprowadzane przez użytkownika modelu).

⁹ Z. Okoń, https://www.linkedin.com/posts/zbigniew-okon_openai-nyt-prawoautorskie-activity-7152621678215278592-jLnF/?originalSubdomain=pl [dostęp 21.05.2024].

¹⁰ *Ibidem*; OpenAI, *OpenAI and journalism*, <https://openai.com/index/openai-and-journalism/> [dostęp 21.05.2024].

¹¹ J. Patterson, A. Gibson, *Deep Learning. Praktyczne wprowadzenie*, Gliwice 2018, s. 19–20.

¹² K. Kunc-Urbańczyk, *Uczenie maszynowe jako pole eksploatacji utworów – 2. Uczenie maszynowe – wyjaśnienie podstawowych pojęć*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, Warszawa 2021.

1.2. Wykorzystanie utworów w procesie tworzenia i funkcjonowania LLM jako pole eksploatacji w postaci zwielokrotniania utworów

Omówienie kwestii wstępnych należy zakończyć na kontekstowym omówieniu zagadnień prawnych. Dla zrozumienia analizy przeprowadzonej w ramach niniejszego tekstu kluczowe jest zrozumienie dwóch pojęć charakterystycznych dla prawa autorskiego, tj. „pola eksploatacji” oraz jednego z jego rodzajów – „zwielokrotniania utworów”.

W polskim porządku prawnym pole eksploatacji utworu to sposób korzystania z tego utworu. Ustawowe przykłady pól eksploatacji zostały wskazane w art. 50 ustawy z dnia 4 lutego 1994 r. o prawie autorskim i prawach pokrewnych¹³ i obejmują poszczególne rodzaje (1) utrwalania i zwielokrotniania utworu, (2) obrotu oryginałem albo egzemplarzami oraz (3) rozpowszechniania utworu w sposób inny niż poprzez obrót utworem. Kluczowe z punktu widzenia niniejszej analizy pole eksploatacji (tj. zwielokrotnianie utworu) oznacza wytwarzanie określoną techniką kolejnego lub kolejnych egzemplarzy utworu (względem pierwszego egzemplarza utworu, stworzonego wskutek utrwalenia) – przy zastosowaniu dowolnych technik (np. druku)¹⁴.

Trzeba jednak zaznaczyć, że o ile uznanie pola eksploatacji za sposób korzystania z utworu można uznać za bezsporne, o tyle w doktrynie toczy się dyskusja co do kryteriów wyodrębniania pól eksploatacji innych niż wymienione w art. 50 pr. aut. – jako że przepis ten zawiera ich katalog otwarty¹⁵. W kontekście analizowanego zagadnienia jest to kwestia o tyle istotna, że wyodrębnienie nowego pola eksploatacji skutkuje koniecznością stosowania do niego przepisów art. 41 ust. 2 pr. aut. (umowy prawnautorskie obejmują pola eksploatacji wyraźnie w nich wymienione) oraz art. 41 ust. 4 pr. aut. (umowa może dotyczyć tylko pól eksploatacji, które są znane w chwili jej zawarcia). W odniesieniu do zarysowanego sporu o wyodrębnianie pól eksploatacji innych niż określone w art. 50 pr. aut. K. Klafkowska-Waśniowska za Elżbietą Traple wskazuje, że pojęcie „pola eksploatacji” wiąże się z formą wykorzystania utworu – odrębną pod względem ekonomicznym lub technicznym¹⁶.

¹³ Ustawa z dnia 4 lutego 1994 r. o prawie autorskim i prawach pokrewnych, Dz. U. z 2022 r. poz. 2509 (dalej: „pr. aut.”).

¹⁴ K. Klafkowska-Waśniowska, [w:] R. Markiewicz (red.), *Ustawy autorskie. Komentarze. Tom I. Komentarz do ustawy o prawie autorskim i prawach pokrewnych*, Warszawa 2021, komentarz do art. 50 pr. aut., teza 25.

¹⁵ K. Gliściński, [w:] A. Michalak (red.), *Ustawa o prawie autorskim i prawach pokrewnych. Komentarz*, Warszawa 2019, komentarz do art. 50 pr. aut., nb 1.

¹⁶ K. Klafkowska-Waśniowska, *op. cit.*, teza 15.

Zdaniem autorów, procesy korzystania z utworów w trakcie tworzenia lub funkcjonowania LLM odbywają się w ramach pola eksploatacji w postaci zwielokrotniania. Należy jednak odnotować, że w doktrynie pojawiły się rozważania co do tego, czy jeden z takich procesów, tj. *machine learning* należy traktować jako swego rodzaju subpole eksploatacji¹⁷. Autorzy przychylają się do tezy, zgodnie z którą, choć z technicznego punktu widzenia możliwe jest uznanie *machine learning'u* za odrębne pole eksploatacji¹⁸, to nie zmienia to jednak faktu, że takie wyodrębnienie nie jest uzasadnione z perspektywy kryterium ekonomicznego. Można zgodzić się ze stwierdzeniem, że *machine learning* pochłania znaczne ilości zasobów i wymaga dużych nakładów finansowych. Należy jednak podkreślić, że *machine learning* jest tylko etapem wytwarzania określonego produktu lub usługi (np. określonego LLM) – autorzy zgadzają się, że dopiero takie produkty lub usługi można oceniać jako posiadające znaczący potencjał ekonomiczny. Sam *machine learning* jako jeden z procesów prowadzących do ich wytworzenia nie ma potencjału wystarczającego, aby uznać go za odrębne pole eksploatacji w rozumieniu art. 50 pr. aut.¹⁹

W konsekwencji poczynionych ustaleń, na potrzeby dalszych rozważań, korzystanie z utworów w ramach omawianych procesów będzie przez autorów postrzegane przez pryzmat pola eksploatacji, jakim jest zwielokrotnianie, dlatego też w dalszej części artykułu przeanalizowane zostanie, kiedy tworzenie i funkcjonowanie LLM prowadzi do zwielokrotniania utworów i czy w danym przypadku zwielokrotnianie to jest objęte dozwolonym użytkowiem. Dodatkowo, na marginesie powyższych rozważań:

- należy jednak stwierdzić, że uznanie jednego z przejawów korzystania z utworów w procesie tworzenia i funkcjonowania LLM za odrębne od zwielokrotniania pole eksploatacji mogłoby wobec takiego odrębnego pola eksploatacji rodzić wątpliwości co do stosowania do niego przepisów o dozwolonym użytku – jak np. art. 4 dyrektywy DSM (więcej szczegółowych rozważań w dalszej części tekstu);

- w orzecznictwie TSUE, bez względu na legislacyjną siatkę terminologiczną korzystania z utworu przyjętą w prawodawstwach krajów członkowskich Unii Europejskiej, niektóre z pól eksploatacji należy uznać za zharmonizowane – w tym

¹⁷ K. Kunc-Urbańczyk, *Uczenie maszynowe jako pole eksploatacji utworów – 6. Możliwość uznania uczenia maszynowego za odrębne subpole zwielokrotniania metodą cyfrową*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *op. cit.*

¹⁸ *Ibidem.*

¹⁹ *Ibidem.*

właśnie prawo do zwielokrotniania²⁰ (obejmujące również utrwalanie²¹) oraz prawo do komunikowania utworu publicznie²². Dzięki temu rozważania poczynione w niniejszej analizie w oparciu o przepisy polskiego porządku prawnego są przynajmniej w pewnym zakresie aktualne w odniesieniu do porządków prawnych pozostałych krajów członkowskich UE.

2. Zwielokrotnianie utworów w związku z tworzeniem LLM

2.1. Przykłady zwielokrotniania utworów w procesie tworzenia LLM

Jak wspomniano, do zwielokrotniania utworów dochodzi na etapie tworzenia LLM. Aby model został wytrenowany na określonym (i ogromnym) zbiorze danych treningowych, dane te muszą zostać najpierw zgromadzone, a następnie wstępnie przetworzone. Z technicznych aspektów czynności gromadzenia danych wynika, że aby możliwe było „nakarmienie” nimi trenowanego LLM, niezbędne jest uprzednie zwielokrotnienie takich danych (a więc również danych stanowiących utwory chronione prawem autorskim) – choćby poprzez zapisanie ich w pamięci komputera.

Pewne wątpliwości może nasuwać to, czy również przetwarzanie zgromadzonych danych, aby uzyskać pożądanej jakości dane treningowe, powinno być kwalifikowane jako zwielokrotnienie utworów. Surowe dane zgromadzone przed rozpoczęciem trenowania LLM (a więc np. pozyskane z różnych źródeł teksty, grafiki, zdjęcia, pliki wideo itp.) muszą przejść tzw. proces wektoryzacji, polegający na konwersji danych na dane liczbowe, tj. do formatu wektorowego umożliwiającego trenowanie algorytmu²³. W procesie wektoryzacji danych dochodzi np. do oczyszczenia tekstów poprzez usunięcie z nich wielkich liter, znaków specjalnych, symboli, cyfr, znaków interpunkcyjnych, spójników oraz do zapisania słów w podstawowej formie gramatycznej²⁴, przetworzenia obrazów poprzez standaryzację wymiarów, zmniejszenie jakości plików, pominięcie jednolitych fragmentów tła oraz ujednolicenie obszarów podobnych barw i normalizację kolorów²⁵. Istnieją wątpliwości, czy

²⁰ Wyrok Trybunału Sprawiedliwości z dnia 29 lipca 2019 r., C-469/17, *Funke Medien*, Dz. Urz. UE C 319/5 z 2019 r.

²¹ K. Klafkowska-Waśniowska, *op. cit.*, teza 23.

²² Wyrok Trybunału Sprawiedliwości z dnia 13 lutego 2014 r., C-466/12, *Svensson*, ZOTSiS 2014/2/I-76.

²³ K. Kunc-Urbańczyk, *Uczenie maszynowe jako pole eksploatacji utworów – 4.2. Wektoryzacja danych*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *op. cit.*

²⁴ *Ibidem*.

²⁵ *Ibidem*.

tego typu operacje na zgromadzonych danych (w tym danych stanowiących utwory) zawsze muszą wiązać się z ponownym zwielokrotnieniem utworów. Zdaniem autorów możliwe są sytuacje, w których w trakcie procesu wektoryzacji danych będzie dochodziło do zwielokrotniania utworów, ale możliwe też są przypadki, w których wektoryzacja danych nie będzie prowadziła do zwielokrotniania – pośrednio wskazuje na to R. Markiewicz, zauważając, że w szerszym kontekście eksploracji tekstów i danych mogą wystąpić przypadki eksploracji, które nie stanowią zwielokrotniania, a taka eksploracja odpowiada „czytaniu utworu” przez człowieka i nie wkracza w monopol prawnoautorski²⁶.

2.2. Zwielokrotnianie utworów w procesie tworzenia LLM a dozwolony użytek utworów

Aby zwielokrotnianie danych (w tym utworów) w ramach szeroko rozumianego etapu tworzenia LLM (a więc obejmującego gromadzenie danych, ich wektoryzację oraz trenowanie modelu) nie naruszało praw podmiotów uprawnionych, twórcy modelu muszą mieć podstawę prawną do takiego korzystania. Podstawę taką mogą stanowić umowy licencyjne zawarte przez twórców LLM z podmiotami uprawnionymi lub przepisy powszechnie obowiązującego prawa. Przedmiotem dalszych rozważań będzie drugi ze wskazanych przypadków.

Wobec tego, że zakres uprawnienia do zwielokrotniania utworów (prawa twórców do zezwalania lub zabrania takiego zwielokrotniania) jest w prawie unijnym (a konsekwentnie także w polskim) ujmowany maksymalnie szeroko²⁷ (por. brzmienie art. 2 dyrektywy 2001/29/WE²⁸), nie ulega wątpliwości, że zwielokrotnianie utworów w związku z tworzeniem LLM co do zasady wkracza w monopol prawnoautorski podmiotów uprawnionych. W doktrynie prawa rozważano, czy podstawą prawną do takiego zwielokrotniania utworów bez zgody podmiotów uprawnionych mogą być przepisy dyrektywy DSM²⁹. Na ten moment, tj. po zmianach w proponowanej treści AI Act (w wersji przyjętej 13 marca 2024 r. przez

²⁶ R. Markiewicz, 2.2. *Eksploracja tekstów i danych (art. 3 i 4)*, [w:] *idem* (red.), *Prawo autorskie na jednolitym rynku cyfrowym. Dyrektywa Parlamentu Europejskiego i Rady (UE) 2019/790*, Warszawa 2021.

²⁷ Z. Okoń, *Wykorzystanie utworów dla potrzeb głębokiego uczenia w świetle europejskiego prawa autorskiego* [w:] K. Flaga-Gieruszyńska, J. Gołaczyński, D. Szostek (red.), *Sztuczna inteligencja, blockchain, cyberbezpieczeństwo oraz dane osobowe. Zagadnienia wybrane*, Warszawa 2019.

²⁸ Dyrektywa 2001/29/WE Parlamentu Europejskiego i Rady z dnia 22 maja 2001 r. w sprawie harmonizacji niektórych aspektów praw autorskich i pokrewnych w społeczeństwie informacyjnym, Dz. Urz. UE L 167 z 22.06.2001 r., s. 10–19 (dalej: „dyrektywa InfoSoc”).

²⁹ Z. Okoń, *op. cit.*

Parlament Europejski i w wersji przyjętej 14 maja 2024 r. przez Radę Europejską), a następnie na gruncie ostatecznie przyjętej wersji tekstu AI Act, zdaniem autorów istnieje możliwość objęcia zwielokrotniania dokonującego się w procesie tworzenia LLM (od etapu gromadzenia danych treningowych do etapu samego trenowania modelu) zakresem dozwolonego użytku utworów z art. 3 i art. 4 dyrektywy DSM w ramach eksploracji tekstów i danych. Wskazuje na to motyw 105 ostatecznej wersji tekstu Aktu o sztucznej inteligencji, który stanowi, że „[...] Rozwój i trenowanie takich modeli wymaga dostępu do ogromnych ilości tekstów, obrazów, materiałów wideo i innych danych. Techniki eksploracji tekstów i danych mogą być w tym kontekście szeroko wykorzystywane do wyszukiwania i analizy takich treści, które mogą być chronione prawem autorskim i prawami pokrewnymi. Każde wykorzystanie treści chronionych prawem autorskim wymaga zezwolenia danego podmiotu praw, chyba że zastosowanie mają odpowiednie wyjątki i ograniczenia dotyczące praw autorskich. Na mocy dyrektywy (UE) 2019/790 wprowadzono wyjątki i ograniczenia umożliwiające, pod pewnymi warunkami, zwielokrotnianie i pobieranie utworów lub innych przedmiotów objętych ochroną do celów eksploracji tekstów i danych. Zgodnie z tymi przepisami podmioty uprawnione mogą zastrzec swoje prawa do swoich utworów lub innych przedmiotów objętych ochroną, aby zapobiec eksploracji tekstów i danych, chyba że odbywa się to do celów badań naukowych. W przypadku gdy prawo do wyłączenia z eksploracji zostało w odpowiedni sposób wyraźnie zastrzeżone, dostawcy LLM muszą uzyskać zezwolenie od podmiotów uprawnionych, jeżeli chcą dokonywać eksploracji tekstów i danych odnośnie do takich utworów”. Jest to również akceptowane przez polską doktrynę prawa, w której wskazuje się, że w ramach eksploracji tekstów i danych mieszczą się takie czynności (etapy) jak:

- 1) identyfikacja odpowiednich danych i uzyskanie do nich dostępu;
- 2) kopiowanie danych w celu ich późniejszej analizy;
- 3) przetworzenie danych do odpowiedniego formatu;
- 4) ekstrakcja tekstów i danych;
- 5) analiza danych w celu identyfikacji związków, zależności, wzorów lub tendencji;
- 6) wykorzystanie rezultatów analizy, tj. użycie danych, przedstawienie rezultatów analizy, weryfikacja wyników (w tym wraz z fragmentami analizowanych tekstów)³⁰.

³⁰ R. Markiewicz, *op. cit.*

Zanim objęcie tworzenia LLM zakresem dozwolonego użytku z dyrektywy DSM zostało szerzej zaakceptowane, w doktrynie rozważano, czy podstawą do zwielokrotniania utworów w takim procesie może być przepis art. 5 dyrektywy InfoSoc. Zgodnie z tym przepisem tymczasowe zwielokrotnianie, mające charakter przejściowy lub dodatkowy, stanowiące integralną i podstawową część procesu technologicznego, którego jedynym celem jest umożliwienie legalnego korzystania z utworu i które nie ma odrębnego znaczenia ekonomicznego, jest wyłączone z zakresu przysługującego twórcy wyłącznego prawa do zwielokrotniania. Za dopuszczalne uznane zostało jednak zwielokrotnianie dokonywane w określonym przebiegu procesu technicznego, tj. pobieranie utworów stanowiących surowe dane, ich przetwarzanie do postaci danych treningowych, a następnie ich przetwarzanie przez sieć neuronową w procesie trenowania i usuwanie po zakończeniu tego procesu, które przebiegało w sposób zautomatyzowany³¹. Rozważania te można jednak uznać obecnie za drugorzędne w związku z przyjęciem dyrektywy DSM oraz brzmieniem ostatecznej wersji tekstu AI Act.

Na marginesie warto wspomnieć, że zarysowane wątpliwości towarzyszyły również polskiemu ustawodawcy, który w 2024 r. podjął kolejną próbę implementacji dyrektywy DSM do polskiego porządku prawnego. W projekcie ustawy z dnia 14 lutego 2024 r.³² proponowane brzmienie art. 26³ ust. 1 pr. aut. zakładało, że nie wolno zwielokrotniać rozpowszechnionych utworów w celu eksploracji tekstów i danych na potrzeby tworzenia generatywnych modeli sztucznej inteligencji, chyba że uprawniony wyraźnie wyraził na to zgodę. Z kolei projekt ustawy przekazany do prac legislacyjnych 16 maja 2024 r.³³ i ostateczna treść uchwalonej ustawy³⁴ dopuszczają możliwość eksploracji tekstów i danych w celu tworzenia generatywnych modeli sztucznej inteligencji – chyba że uprawniony wyraźnie sprzeciwi się takiemu korzystaniu z jego utworów. Zmiana brzmienia propozycji i ostatecznej wersji polskich przepisów implementujących dyrektywę DSM jest w naszej ocenie konsekwencją wskazanych wyżej zmian, jakim podlegały kolejne wersje projektu AI Act.

³¹ Z. Okoń, *op. cit.*

³² Projekt ustawy z dnia 14 lutego 2024 r. o zmianie ustawy o prawie autorskim i prawach pokrewnych oraz niektórych innych ustaw, <https://legislacja.rcl.gov.pl/docs/2/12382002/13037394/13037395/dokument656778.pdf> [dostęp 22.05.2024].

³³ Rządowy projekt ustawy o zmianie ustawy o prawie autorskim i prawach pokrewnych oraz niektórych innych ustaw, wniesiony do Sejmu 16 maja 2024, <https://www.sejm.gov.pl/sejm10.nsf/agent.xsp?symbol=RPL&Id=RM-0610-36-24> [dostęp 22.05.2024].

³⁴ Ustawa z dnia 26 lipca 2024 r. o zmianie ustawy o prawie autorskim i prawach pokrewnych, ustawy o ochronie baz danych oraz ustawy o zbiorowym zarządzaniu prawami autorskimi i prawami pokrewnymi, Dz. U. z 2024 r. poz. 1254.

3. Zwielokrotnianie w ramach zjawiska zapamiętywania

3.1. Przykłady zwielokrotniania w ramach zjawiska zapamiętywania

Jak wskazują Ivo Emanuilov i Thomas Margoni, powołując się na badania przeprowadzone przez naukowców zrzeszonych w USENIX Association, zostało empirycznie wykazane, że LLM mogą wykazywać zachowanie podobne do tego, co w odniesieniu do pracy ludzkiego mózgu nazywamy pamięcią fotograficzną lub ejdetyczną, tj. zdolnością do przywoływania informacji po ich jednokrotnym obejrzeniu³⁵. Należy w tym miejscu podkreślić za ww. autorami, że nie każdy algorytm uczenia maszynowego zachowuje się w taki sam sposób i istnieją przypadki, w których zapamiętywanie jest zamierzoną cechą procesu uczenia modelu³⁶. Przedmiotem niniejszej analizy jest jednak takie zjawisko zapamiętywania, które stanowi niezamierzoną cechę tego procesu. Innymi słowy, przedmiotem analizy jest zapamiętywanie przez LLM danych treningowych niezbędne do tego, aby LLM dokonywały bardziej efektywnego uogólniania (tj. zdolności LLM do adaptacji i prawidłowego reagowania na nieznane mu wcześniej dane, które nie wchodziły w zakres danych treningowych)³⁷.

W związku z poczynionym spostrzeżeniem, że wynikający z art. 2 dyrektywy InfoSoc zakres prawa twórców do zezwalania lub zabrania zwielokrotniania utworów jest w prawie unijnym ujmowany szeroko, należy uznać, że zwielokrotnianie utworów w toku procesu zapamiętywania wkracza w monopol prawnoautorski. Po raz kolejny pojawia się zatem pytanie, czy takie zwielokrotnianie może odbywać się w granicach dozwolonego użytku wynikającego z przepisów prawa.

3.2. Zwielokrotnianie utworów w ramach zjawiska zapamiętywania a dozwolony użytek utworów

W pierwszej kolejności należy zastanowić się, czy zwielokrotnianie utworów w procesie zapamiętywania jest objęte zakresem dozwolonego użytku wynikającego z dyrektywy DSM. Jak wynika z przywołanych wcześniej art. 3 i art. 4 dyrektywy DSM, przepisy te obejmują zwielokrotnianie i pobieranie utworów do celów

³⁵ I. Emanuilov, T. Margoni, *Forget me not: memorisation in generative sequence models trained on open source licensed code*, 27.02.2024, s. 11, <https://zenodo.org/records/10635479>; [dostęp 20.05.2024].

³⁶ *Ibidem*.

³⁷ L. Goetz, N. Seedat, R. Vandersluis, *Generalization – a key challenge for responsible AI in patient-facing clinical applications*, „Nature. npj Digital Medicine” 2024, 7, 126, <https://www.nature.com/articles/s41746-024-01127-3> [dostęp 22.05.2024].

eksploracji tekstów i danych. Do rozstrzygnięcia pozostaje zatem, czy zwielokrotnianie utworów w procesie ich zapamiętywania przez LLM (ale przed ich zwielokrotnieniem w formie outputu – o czym niżej) stanowi zwielokrotnienie mieszczące się w sformułowaniu „do celów eksploracji tekstów i danych”.

Zdaniem autorów, zwielokrotnienie utworów odbywające się od momentu gromadzenia surowych danych treningowych, przez proces wektoryzacji danych, aż do momentu zakończenia trenowania LLM (na określonych danych – utworach) wchodzi w zakres dozwolonego użytku związanego z eksploracją tekstów i danych. Autorzy twierdzą jednak, że zwielokrotnianie utworów odbywające się przez to, że wytrenowany LLM zapamiętał określone elementy danych treningowych, takim dozwolonym użytkowaniem nie jest już objęte. Tezę tę autorzy opierają przede wszystkim na założeniu, że dozwolone zwielokrotnianie wynikające z art. 3 i art. 4 dyrektywy DSM stanowi wyjątek, który zgodnie z zasadą *exceptiones non sunt extendendae* nie powinien być wykładany rozszerzająco. Takim nadużyciem byłoby z kolei stwierdzenie, że zwielokrotnianie odbywające się faktycznie już po procesie konwersji danych treningowych do formatu odczytywalnego przez sieć neuronową i po procesie trenowania takiej sieci nadal jest objęte zakresem dozwolonego użytku wynikającego z omawianych przepisów.

Podobne spostrzeżenie należy poczynić względem „wyrwy” tworzonej w monopolu prawnoautorskim przez art. 5 ust. 1 lit. b) dyrektywy InfoSoc. Jak wspomniano, przewiduje on, że tymczasowe zwielokrotnianie, mające charakter przejściowy lub dodatkowy, które stanowi integralną i podstawową część procesu technologicznego i którego jedynym celem jest umożliwienie legalnego korzystania z utworu i które nie ma odrębnego znaczenia ekonomicznego, jest wyłączone z zakresu przysługującego twórcy wyłącznego prawa do zwielokrotniania. Nie sposób jednak uznać, że zwielokrotnianie odbywające się wskutek zapamiętania danych treningowych przez LLM ma charakter przejściowy lub dodatkowy. Ma to związek ze stwierdzeniem, że zapamiętywanie mniejszych lub większych fragmentów danych treningowych po pierwsze – może stanowić trwałe zakodowanie danych w LLM i po drugie – ma wpływ na zwiększenie efektywności zdolności uogólniania, tj. na zmniejszenie niedokładności modelu w przewidywaniu wartości wyników na nieznanym mu dotąd danych. Konsekwentnie nie można stwierdzić, że zwielokrotnianie nietymczasowe i nieobjęjące automatycznego usuwania utworów nie jest objęte zakresem dozwolonego użytku przewidzianym w art. 5 ust. 1 lit. b) dyrektywy InfoSoc³⁸.

³⁸ R. Markiewicz, *ChatGPT i prawo autorskie Unii Europejskiej*, „Zeszyty Naukowe Uniwersytetu Jagiellońskiego” 2023, nr 2(160), s. 152.

4. Zwielokrotnianie w ramach zjawiska regurgitacji treści przez LLM

4.1. Przykłady zwielokrotniania w ramach zjawiska regurgitacji treści przez LLM

Konsekwencją opisanego zjawiska zapamiętywania utworów wchodzących w skład danych treningowych są przypadki, gdy po wprowadzeniu odpowiedniego promptu lub promptów generacja tworzona przez LLM zawiera całe fragmenty danych, na których model został wytrenowany.

Na gruncie sporu „New York Times” przeciwko OpenAI, zjawisko takie zostało przez pozwanego dostawcę ChataGPT³⁹ nazwane „regurgitacją”⁴⁰ – znanym z medycyny określeniem zjawiska niekontrolowanego cofania się treści pokarmowej z żołądka do przełyku⁴¹. W toku sporu „New York Times” wykazuje, że empirycznie zweryfikował, że po wprowadzeniu do promptu początkowych fragmentów swoich artykułów uzyskiwał w outpucie generowanym przez ChatGPT praktycznie identyczne fragmenty swoich tekstów. Według wydawcy oznacza to, że – po pierwsze – LLM stworzony przez OpenAI został wytrenowany na tekstach, do których prawa autorskie przysługują gazecie i – po drugie – że możliwe są przypadki zwielokrotniania tych tekstów w treściach generowanych przez sztuczną inteligencję⁴².

4.2. Zwielokrotnianie utworów w ramach zjawiska regurgitacji treści przez LLM

Zdaniem autorów przykłady zwielokrotniania utworów wskazane w pkt 5.1. pozostawiają najmniej wątpliwości co do ich nieuprawnionego charakteru (zakładając oczywiście, że w danym przypadku nie udzielono odpowiedniej licencji). Z pewnością bowiem takie zwielokrotnianie nie może być objęte zakresem dozwolonego użytku, czy to z art. 3 lub art. 4 dyrektywy DSM (generowanie treści nie stanowi bo-

³⁹ R. Metz, *OpenAI Claims New York Times Isn't 'Telling the Full Story' in Suit*, Bloomberg, 8.01.2024, <https://www.bloomberg.com/news/articles/2024-01-08/openai-says-new-york-times-not-telling-the-full-story-in-suit?embedded-checkout=true> [dostęp 22.05.2024].

⁴⁰ Tak też w referacie *Naruszenie prawa autorskiego w wytworach generatywnej sztucznej inteligencji*, ogłoszonym przez Z. Okonia 22.03.2024 r. podczas 10. Forum Prawa Mediów Elektronicznych.

⁴¹ J. Ryżko, *Gastroenterologia pediatryczna – pytania i odpowiedzi*, „Gastroenterologia Praktyczna” 2012, nr 1, s. 81–82.

⁴² M. Bastian, *OpenAI claims New York Times' prompting strategy violates its terms of service*, 6.01.2024, <https://the-decoder.com/openai-claims-new-york-times-prompting-strategy-violates-its-terms-of-service/> [dostęp 22.05.2024].

wiem eksploracji tekstów i danych), czy to z art. 5 ust. 1 lit. b) dyrektywy InfoSoc – nie jest to wszakże zwielokrotnianie, które ma charakter przejściowy lub dodatkowy.

Na dodatek generowanie treści przez LLM wskutek tzw. regurgitacji, nawet jeśli stanowiłoby dozwolone zwielokrotnianie (o czym, jak wspomniano, w opinii autorów nie może być mowy), to i tak stanowiłoby to wykraczające poza ramy dozwolonego użytku (zarówno z dyrektywy DSM, jak i z dyrektywy InfoSoc) korzystanie z utworu poprzez ich publiczne udostępnianie⁴³.

5. Wnioski

Podsumowując, rozpatrywanie zagadnienia zwielokrotniania utworów w procesach tworzenia i funkcjonowania LLM prowadzi do następujących wniosków:

- do zwielokrotniania danych stanowiących utwory może dochodzić na różnych etapach tych procesów, tj. (1) podczas zbierania danych treningowych, (2) w trakcie procesu wektoryzacji danych, (3) w toku dalszego trenowania LLM, (4) wskutek zapamiętywania utworów zawartych w danych treningowych, aż po (4) możliwą regurgitację utworów w treściach generowanych przez LLM;

- nie każdy z ww. przypadków zwielokrotniania jest objęty zakresem dozwolonego użytku, a więc potencjalnie może naruszać prawa podmiotów uprawnionych;

- zdaniem autorów zwielokrotnienie utworów odbywające się od momentu gromadzenia surowych danych treningowych, przez proces wektoryzacji danych, aż do momentu zakończenia trenowania LLM na określonych danych (utworach) wchodzi w zakres dozwolonego użytku związanego z eksploracją tekstów i danych;

- w ocenie autorów zwielokrotnianie utworów odbywające się przez to, że wytrenowany LLM zapamiętał określone elementy danych treningowych, nie jest objęte dozwolonym użytkowaniem wynikającym z art. 3 i art. 4 dyrektywy DSM (i konsekwentnie – art. 26² oraz art. 26³ polskiej ustawy o prawie autorskim i prawach pokrewnych znowelizowanej w lipcu 2024 r. ustawą implementującą dyrektywę DSM do polskiego porządku prawnego⁴⁴). Przepisy te ustanawiają wyjątek od monopolu prawnoautorskiego, który nie powinien być wykładany rozszerzająco;

- zdaniem autorów zwielokrotnianie danych stanowiących utwory, odbywające się wskutek zapamiętania danych treningowych przez LLM, nie ma charakteru

⁴³ I. Emanuilov, T. Margoni, *Memorisation...*

⁴⁴ Ustawa z dnia 26 lipca 2024 r. o zmianie ustawy o prawie autorskim i prawach pokrewnych, ustawy o ochronie baz danych oraz ustawy o zbiorowym zarządzaniu prawami autorskimi i prawami pokrewnymi, Dz. U. z 2024 r. poz. 1254.

przejściowego lub dodatkowego, a także nie obejmuje automatycznego usuwania utworów – dlatego też nie jest objęte zakresem dozwolonego użytku przewidzianym w art. 5 ust. 1 lit. b) dyrektywy InfoSoc;

– w ocenie autorów zwielokrotnianie utworów w ramach zjawiska regurgitacji treści przez LLM nie może być objęte zakresem dozwolonego użytku ani z art. 3 lub art. 4 dyrektywy DSM (generowanie treści nie stanowi bowiem eksploracji tekstów i danych), ani z art. 5 ust. 1 lit. b) dyrektywy InfoSoc (nie jest to zwielokrotnianie, które ma charakter przejściowy lub dodatkowy).

Należy również dodać, że wnioski płynące z analizy mają charakter wyjątkowo praktyczny – choćby z uwagi na znaczną liczbę sporów sądowych, w których dostawcy LLM są pozywani przez kolejnych twórców (por. przywołany w tekście przykład sporu „New York Times” przeciwko OpenAI). Autorzy pragną jednak przypomnieć, że są to spory sądowe odbywające się na gruncie prawa amerykańskiego i w tradycji prawa *common law*. Konsekwentnie, ewentualne rozstrzygnięcia nie wpłyną na wykładnię omawianych w tekście przepisów dyrektywy DSM i dyrektywy InfoSoc. Z takiego założenia wychodzą również najprawdopodobniej dostawcy LLM, którzy nie chcą ryzykować zmniejszenia popytu na ich usługi (do czego mogłoby dojść, jeśli użytkownicy narzędzi AI obawialiby się, że korzystanie przez nich z tych narzędzi mogłoby uzasadniać kierowanie przeciwko nim roszczeń z tytułu praw autorskich), wprowadzając w swoich warunkach umów tzw. klauzule indemnifikacyjne, w których, pod pewnymi warunkami, zobowiązują się zwolnić użytkownika z odpowiedzialności zarówno za ewentualne naruszenia cudzych praw w związku z danymi wykorzystanymi do trenowania ich LLM, jak i w związku z naruszeniami związanymi z generacją treści przez AI⁴⁵.

Bibliografia

Akty prawne

Dyrektywa 2001/29/WE Parlamentu Europejskiego i Rady z dnia 22 maja 2001 r. w sprawie harmonizacji niektórych aspektów praw autorskich i pokrewnych w społeczeństwie informacyjnym (Dz. Urz. UE L 167 z 22.06.2001 r., s. 10–19).

⁴⁵ Por. np. Copilot Copyright Commitment firmy Microsoft, <https://blogs.microsoft.com/on-the-issues/2023/09/07/copilot-copyright-commitment-ai-legal-concerns/> [dostęp 22.05.2024] lub Szczegółowe warunki korzystania z usługi Google Workspace, <https://workspace.google.com/terms/service-terms/> [dostęp 22.05.2024].

Dyrektywa Parlamentu Europejskiego i Rady (UE) 2019/790 z dnia 17 kwietnia 2019 r. w sprawie prawa autorskiego i praw pokrewnych na jednolitym rynku cyfrowym oraz zmiany dyrektyw 96/9/WE i 2001/29/WE (Dz. Urz. UE L 130 z 17.05.2019 r., s. 92–125).

Projekt ustawy z dnia 14 lutego 2024 r. o zmianie ustawy o prawie autorskim i prawach pokrewnych oraz niektórych innych ustaw z dnia 14 lutego 2024 r., <https://legislacja.rcl.gov.pl/docs//2/12382002/13037394/13037395/dokument656778.pdf> [dostęp 22.05.2024 r.].

Propozycja Rozporządzenia Parlamentu Europejskiego i Rady 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), tekst w wersji przyjętej w dniu 14 maja 2024 r., 2021/0106(COD), PE-CONS 24/24, <https://data.consilium.europa.eu/doc/document/PE-24-2024-INIT/pl/pdf> [dostęp 21.05.2024].

Rezolucja ustawodawcza Parlamentu Europejskiego z dnia 13 marca 2024 r. w sprawie 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)0206 – C9-0146/2021 – 2021/0106(COD)), https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_PL.pdf [dostęp 17.05.2024].

Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 z dnia 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 z 12.07.2024 r.).

Rządowy projekt ustawy o zmianie ustawy o prawie autorskim i prawach pokrewnych oraz niektórych innych ustaw, wniesiony do Sejmu w dniu 16 maja 2024 r., <https://www.sejm.gov.pl/sejm10.nsf/agent.xsp?symbol=RPL&Id=RM-0610-36-24>, [dostęp 22.05.2024].

Ustawa z dnia 26 lipca 2024 r. o zmianie ustawy o prawie autorskim i prawach pokrewnych, ustawy o ochronie baz danych oraz ustawy o zbiorowym zarządzaniu prawami autorskimi i prawami pokrewnymi (Dz. U. z 2024 r. poz. 1254).

Ustawa z dnia 4 lutego 1994 r. o prawie autorskim i prawach pokrewnych (Dz. U. z 2022 r. poz. 2509, t.j. z dnia 6.12.2022 r.).

Orzecznictwo

Wyrok Trybunału Sprawiedliwości z dnia 13 lutego 2014 r., C-466/12, *Svensson*, ZOTSiS 2014/2/I-76

Wyrok Trybunału Sprawiedliwości z dnia 29 lipca 2019 r., C-469/17, *Funkke Medien* (Dz. Urz. UE C 319/5 z 23.09.2019 r.).

Źródła

Bastian M., *OpenAI claims New York Times' prompting strategy violates its terms of service*, 6.01.2024, <https://the-decoder.com/openai-claims-new-york-times-prompting-strategy-violates-its-terms-of-service/> [dostęp 22.05.2024].

Carlini N, Tramèr F., Wallace E., Jagielski M., Herbert-Voss A., Lee K., Roberts A., Brown T., Erlingsson D.U., Oprea A., Raffel C, *Extracting Training Data from Large Language Models*, <https://www.usenix.org/conference/usenixsecurity21/presentation/carlini-extracting> [dostęp 21.05.2024].

- Copilot Copyright Commitment firmy Microsoft, <https://blogs.microsoft.com/on-the-issues/2023/09/07/copilot-copyright-commitment-ai-legal-concerns/> [dostęp 22.05.2024].
- Emanuilov I., Margoni T., *Forget me not: memorisation in generative sequence models trained on open source licensed code*, 27.02.2024, <https://zenodo.org/records/10635479> [dostęp 20.05.2024].
- Emanuilov I., Margoni T., *Memorisation in generative models and EU copyright law: an interdisciplinary view*, Kluwer Copyright Blog, 26.03.2024, <https://copyrightblog.kluweriplaw.com/2024/03/26/memorisation-in-generative-models-and-eu-copyright-law-an-interdisciplinary-view/> [dostęp 21.05.2024].
- Goetz L., Seedat N., Vandersluis R., *Generalization – a key challenge for responsible AI in patient-facing clinical applications*, „Nature. npj Digital Medicine” 2024, 7, 126, <https://www.nature.com/articles/s41746-024-01127-3> [dostęp 22.05.2024].
- Metz R., *OpenAI Claims New York Times Isn't 'Telling the Full Story' in Suit*, Bloomberg, 8.01.2024, <https://www.bloomberg.com/news/articles/2024-01-08/openai-says-new-york-times-not-telling-the-full-story-in-suit?embedded-checkout=true>, [dostęp: 22.05.2024].
- O’Keefe D., *How Does AI Model Training Work?*, appian, 24.01.2024, <https://appian.com/blog/acp/ai/how-does-ai-model-training-work.html> [dostęp 21.05.2024].
- Okoń Z., https://www.linkedin.com/posts/zbigniew-okon_openai-nyt-prawoautorskie-activity-7152621678215278592-jLnF/?originalSubdomain=pl [dostęp 21.05.2024].
- Okoń Z., *Naruszenie prawa autorskiego w wytworach generatywnej sztucznej inteligencji*, referat wygłoszony podczas 10. Forum Prawa Mediów Elektronicznych.
- Szczegółowe warunki korzystania z usługi Google Workspace, <https://workspace.google.com/terms/service-terms/> [dostęp 22.05.2024].
- Tan J., *What Is AI Model Training?*, 30.04.2024, <https://engage-ai.co/what-ai-model-training/> [dostęp 21.05.2024].

Literatura

- Gliściński K., [w:] A. Michalak (red.), *Ustawa o prawie autorskim i prawach pokrewnych. Komentarz*, Warszawa 2019, komentarz do art. 50 pr. aut.
- Klaflowska-Waśniowska K., [w:] R. Markiewicz (red.), *Ustawy autorskie. Komentarze. Tom I. Komentarz do ustawy o prawie autorskim i prawach pokrewnych*, Warszawa 2021, komentarz do art. 50 pr. aut.
- Kunc-Urbańczyk K., *Uczenie maszynowe jako pole eksploatacji utworów – 2. Uczenie maszynowe – wyjaśnienie podstawowych pojęć*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, Warszawa 2021.
- Kunc-Urbańczyk K., *Uczenie maszynowe jako pole eksploatacji utworów – 4.2. Wektoryzacja danych*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, Warszawa 2021.
- Kunc-Urbańczyk K., *Uczenie maszynowe jako pole eksploatacji utworów – 6. Możliwość uznania uczenia maszynowego za odrębne subpole zwielokrotniania metodą cyfrową*, [w:] B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, Warszawa 2021.
- Markiewicz R., 2.2. *Eksploracja tekstów i danych (art. 3 i 4)*, [w:] R. Markiewicz (red.), *Prawo autorskie na jednolitym rynku cyfrowym. Dyrektywa Parlamentu Europejskiego i Rady (UE) 2019/790*, Warszawa 2021.

- Markiewicz R., *ChatGPT i prawo autorskie Unii Europejskiej*, „Zeszyty Naukowe Uniwersytetu Jagiellońskiego” 2023, nr 2(160).
- Okoń Z., *Wykorzystanie utworów dla potrzeb głębokiego uczenia w świetle europejskiego prawa autorskiego*, [w:] K. Flaga-Gieruszyńska, J. Gołaczyński, D. Szostek (red.), *Sztuczna inteligencja, blockchain, cyberbezpieczeństwo oraz dane osobowe. Zagadnienia wybrane*, Warszawa 2019.
- Patterson J., Gibson A., *Deep Learning. Praktyczne wprowadzenie*, Gliwice 2018.
- Ryżko J., *Gastroenterologia pediatryczna – pytania i odpowiedzi*, „Gastroenterologia Praktyczna” 2012, nr 1.

Reproduction of works in the processes of creation and functioning of large language models

Abstract

With regard to so-called large language models, one of the main issues is the threat to the exclusive rights of creators – and in particular – copyright. The aforementioned threat materialises both at the stage of the creation of the LLM and at the stage of the subsequent functioning of such model. On the subject of the use of others' works, the disputes in the legal doctrine are mainly related to the scope of fair use of works within the framework of text and data mining, to the so-called phenomenon of memorisation of training data by LLM and to the phenomenon of so-called regurgitation, i.e. the “return” of fragments of protected training data by LLM in generated content. The disputes indicated were fuelled by the procedure of the latest AI legislation, i.e. primarily the Artificial Intelligence Act, of which subsequent versions appeared in Q1 2024 – and the latest version of which was adopted on 13th of June 2024 – including references to provisions on text and data mining under the Digital Single Market Copyright Directive. The purpose of this text is to analyse the different stages of the creation and operation of LLM and to verify the reproduction of protected works that takes place during these stages in the context of the possible application of the fair use provisions.

Keywords

AI, work, exploitation, fair use, text, data, mining, TDM