

SZTUCZNA INTELIGENCJA

w kontekście ochrony danych osobowych

materiały pokonferencyjne



Autorzy materiałów:

Agnieszka Rapcewicz

Akademia Leona Koźmińskiego, Stowarzyszenie.AI, Grupa ANG S.A.

Dominik Lubasz

Lubasz i Wspólnicy Kancelaria Radców Prawnych sp. k.

Maciej Gawroński, Maciej Grzesiuk

Kancelaria Gawroński & Partners

Michał Lewandowski

Cloudtechnologies.pl

Redaktor prowadzący:

Mirosław Sanek

Tomasz Soczyński

Opracowanie redakcyjne:

Natalia Misiuk

Balbina Hermanowicz

Autorstwo poszczególnych rozdziałów:

Agnieszka Rapcewicz

Akademia Leona Koźmińskiego, Stowarzyszenie.AI, Grupa ANG S.A.

Dominik Lubasz

Lubasz i Wspólnicy – Kancelaria Radców Prawnych sp. k.

Maciej Gawroński, Maciej Grzesiuk

Kancelaria Gawroński & Partners

Michał Lewandowski

Cloudtechnologies.pl

Urząd Ochrony Danych Osobowych

ul. Stawki 2,

00-193 Warszawa



SPIS TREŚCI

WSTĘP.....	2
Sztuczna inteligencja w procesie oceny zdolności kredytowej.....	4
Autor: adw. Agnieszka Rapcewicz	
Konstruowanie humanocentrycznej SI z wykorzystaniem instrumentów prawnych ochrony danych osobowych.....	19
Autor: dr Dominik Lubasz	
Jaki plan dla Sztucznej Inteligencji ma Unia Europejska?.....;	30
Autorzy: r.pr. Maciej Gawroński, Maciej Grzesiuk	
Obiektywnie czy subiektywnie? Jak patrzeć na dane?.....	46
Autor: Michał Lewandowski	

WSTĘP

Prezentujemy Państwu materiały, które stanowią podsumowanie panelu dyskusyjnego pt. „Sztuczna inteligencja – w kontekście ochrony danych osobowych”. Wydarzenie było doskonałą okazją, aby przyjrzeć się w większym stopniu korzyściom, ale również zagrożeniom, jakie wiążą się z przetwarzaniem danych osobowych przez systemy komputerowe symulujące ludzkie myślenie.

Postęp związany z innowacjami musi odbywać się z uwzględnieniem zasad określonych w RODO, w tym minimalizacji danych, prawidłowości czy integralności. W tym celu niezbędne jest podjęcie odpowiednich działań i wypracowanie takich rozwiązań, które przyczynią się do zwiększenia bezpieczeństwa oraz świadomości na temat zagrożeń wynikających z korzystania z dobrodziejstw nowych technologii.

W materiałach omówione zostały tematy związane z wykorzystywaniem sztucznej inteligencji w różnych branżach oraz związanych z tym regulacji w zakresie danych osobowych i prywatności. Poruszono również kwestie konstruowania humanocentrycznej SI z wykorzystaniem instrumentów prawnych ochrony danych osobowych.

Szanowni Państwo, Drodzy Czytelnicy,

oddaję w Państwa ręce publikację, będącą rezultatem pracy uczestników panelu dyskusyjnego „Sztuczna Inteligencja – w kontekście ochrony danych osobowych”, który odbył się 26 kwietnia 2021 r.

Niewątpliwie dynamiczny rozwój różnych form działalności przy użyciu sztucznej inteligencji wiąże się z wykorzystaniem danych, w tym również danych osobowych. Możliwość gromadzenia i przetwarzania tak dużej ilości informacji, może prowadzić do naruszenia prywatności innych osób. Dlatego Urząd Ochrony Danych Osobowych, dużą wagę przywiązuje do kwestii dotyczących odpowiednich rozwiązań prawnych regulujących zasady ochrony danych osobowych pozyskiwanych za jej pomocą. Widząc duże zainteresowanie tą tematyką od lat staramy się zwiększać świadomość społeczeństwa na temat zagrożeń, konsekwencji, zabezpieczeń i praw przysługujących każdej osobie w związku z przetwarzaniem jej danych osobowych.

Dziękuję wszystkim osobom, które przyczyniły się do powstania materiałów pokonferencyjnych „Sztuczna inteligencja – w kontekście ochrony danych osobowych”. Szczególne podziękowania kieruję w stronę Autorów, za przygotowanie wartościowych artykułów i zaangażowanie w powstanie publikacji.

Jestem przekonany, że publikacja ta będzie dla Państwa cennym źródłem wiedzy i porad przydatnych w życiu codziennym w zderzeniu ze sztuczną inteligencją w kontekście ochrony danych i prywatności.

Jan Nowak

Prezes Urzędu Ochrony Danych Osobowych

Sztuczna inteligencja w procesie oceny zdolności kredytowej

Autor: adw. Agnieszka Rapcewicz

Akademia Leona Koźmińskiego, Stowarzyszenie.AI, Grupa ANG S.A.

1. Wprowadzenie – pojęcie sztucznej inteligencji

Pojęcie sztucznej inteligencji może być różnie rozumiane i wykorzystywane do opisanego wielu rozwiązań i mechanizmów. Aktualnie brak normatywnej definicji sztucznej inteligencji. Na wstępie zasadne jest więc wyjaśnienie, jak powyższe pojęcie będzie rozumiane w ramach niniejszego opracowania.

Sztuczną inteligencję można rozumieć ściśle jako system (maszynę) naśladujący funkcje poznawcze człowieka (rozumienie, wyrażanie, postrzeganie, obliczanie, zapamiętywanie, organizowanie, rozumowanie, wyobrażanie, tworzenie i rozwiązywanie problemów). Aktualnie jednak pojęcie sztucznej inteligencji jest powszechnie używane przez praktyków i teoretyków prawa oraz dziedzin technicznych w odniesieniu do uczenia maszynowego (machine learning) i algorytmów głębokiego uczenia się (deep learning). Są to jedynie pewne techniki czy narzędzia wykorzystywane przy tworzeniu mechanizmów mających na celu analizę informacji i zdarzeń, w celu automatyzacji podejmowania decyzji i wykonywania określonych działań.

Zdefiniowania pojęcia sztucznej inteligencji podejmują się różne organizacje międzynarodowe czy instytucje krajowe, przygotowując wytyczne lub rekomendacje w zakresie wykorzystywania systemów sztucznej inteligencji. Przykładem jest definicja

zawarta w Rekomendacjach Rady OECD dot. Sztucznej Inteligencji[1]. Zgodnie z tym dokumentem sztuczna inteligencja to system oparty na koncepcji maszyny, która może wpływać na środowisko rzeczywiste lub wirtualne, formułując zalecenia, przewidywania lub decyzje dotyczące zadanego przez człowieka zestawu celów. Czyni to, wykorzystując dane wejściowe, dane maszynowe lub ludzkie do postrzegania rzeczywistych lub wirtualnych środowisk, streszczania takiego postrzegania w modele ręcznie lub automatycznie, wykorzystywania interpretacji modeli do formułowania opcji wyników. Wskazana definicja sformułowana przez OECD została przyjęta za własną w Polityce dla rozwoju sztucznej inteligencji w Polsce od roku 2020 ustanowionej uchwałą Rady Ministrów z dnia 28 grudnia 2020 r.[2]

Definicja normatywna sztucznej inteligencji została zaproponowana w projekcie rozporządzenia Parlamentu i Rady w sprawie sztucznej inteligencji[3]. Zgodnie z nim, a konkretnie art. 3 pkt (1), system sztucznej inteligencji (system SI) oznacza oprogramowanie rozwijane poprzez jedną lub więcej technik i metod wskazanych w załączniku I, które może, dla zdefiniowanych przez człowieka celów, generować wyniki takie jak treść, predykcje, rekomendacje, czy decyzje wpływające na środowiska, z którymi wchodzi w interakcje[4]. We wspomnianym załączniku I wymieniono następujące systemy sztucznej inteligencji: metody uczenia maszynowego (machine learning), w tym nadzorowanego, nienadzorowanego i uczenia się przez wzmocnienie, z wykorzystaniem szerokiej gamy metod, w tym głębokiego uczenia się; metody oparte na logice i wiedzy, w tym reprezentacja wiedzy, programowanie indukcyjne (logiczne), bazy wiedzy, silniki wnioskowania i dedukcji, rozumowanie (symboliczne) i systemy eksperckie; metody statystyczne, estymacja bayesowska, metody poszukiwawcze i optymalizacyjne.

Jak wynika z wyżej przytoczonych dokumentów, generalnie przyjmowana jest szeroka definicja sztucznej inteligencji. Tak również jest ona rozumiana na potrzeby niniejszego opracowania, tj. jako mechanizmy i techniki analizujące (oparte na wykorzystaniu algorytmów), interpretujące zdarzenia i informacje (dane), w celu automatyzacji podejmowania decyzji i wykonywania określonych działań. Możliwe wyniki wspomnianej

analizy, w zależności od modelu, to co do zasady predykcja (np. dana osoba będzie spłacać pożyczkę, posiada zdolność kredytową), rekomendacja (np. określony produkt będzie interesujący dla danego klienta, można wyświetlić mu daną reklamę) lub klasyfikacja (np. wiadomość zostanie zakwalifikowana jako spam). Co więcej, warto podkreślić, że wykorzystywanie sztucznej inteligencji do osiągnięcia wspomnianych wyników nie musi polegać na podejmowaniu decyzji bez jakiegokolwiek udziału człowieka (tj. gdy wynik dostarczony przez SI stanowi samodzielną podstawę decyzji). Stosowanie systemów sztucznej inteligencji może być jedynie elementem szerszego procesu decyzyjnego, w którym analiza dostarczona przez system podlega jeszcze weryfikacji lub uzupełnieniu człowieka.

2. Sztuczna inteligencja w sektorze finansowym

Sztuczna inteligencja może mieć zastosowanie wszędzie tam, gdzie istnieje potrzeba przeanalizowania dużej liczby danych i wyciągnięcia z nich określonych wniosków. Stosowanie systemów sztucznej inteligencji służy przede wszystkim przyspieszeniu i usprawnieniu określonych procesów. Można wyróżnić rozwiązania, które mogą być stosowane w różnych sektorach, jak i takie, które są charakterystyczne w szczególności dla omawianej branży.

Mówiąc o sektorze finansowym można wskazać systemy sztucznej inteligencji wykorzystywane do:

- automatyzacji istotnych procesów, zazwyczaj takich, które wymagają analizy dużej ilości danych;
- usprawnienia komunikacji z klientami na początkowym etapie, do czego mogą być wykorzystywane chatboty lub voiceboty;
- tworzenia profili klientów i rekomendacji w zakresie produktów/usług, jakie można im zaoferować;

- tworzenia prognoz dla danej instytucji i jej klientów, na podstawie których planowane są przyszłe działania;
- oceny zdolności kredytowej;
- analizowania transakcji pod kątem wykrywania oszustw.

Powyższe działania wymagają dokonania oceny dużej ilości danych, wśród których mogą znajdować się dane osobowe. Warto pamiętać o tym, że wykorzystywanie systemów sztucznej inteligencji może wiązać się z przetwarzaniem danych osobowych praktycznie na każdym etapie wykorzystywania omawianych rozwiązań, tj. na etapie trenowania, testowania i wdrożenia. Co więcej, z danymi osobowymi możemy mieć do czynienia nie tylko w zakresie danych wsadowych (tj. dostarczonych systemowi i poddawanych analizie), ale również w zakresie danych wynikowych (wynik analizy stanowi nowe dane osobowe wytworzone przez system sztucznej inteligencji). Przykładem działań, w których przetwarzane są dane osobowe zarówno w zakresie informacji wejściowych, jak i wynikowych, jest analiza zdolności kredytowej. Bank lub inna instytucja na podstawie informacji m.in. o wieku, stanie cywilnym, osobach na utrzymaniu klienta, jego wykształceniu, zatrudnieniu, uzyskiwanych dochodach, źródle utrzymania, posiadanych zobowiązaniach itd. wyciąga wnioski o możliwości spłaty przez niego zaciągniętego kredytu wraz z odsetkami w umówionym terminie (informacja ta stanowi dane osobowe wytworzone w procesie oceny zdolności kredytowej).

3. Sztuczna inteligencja i przetwarzanie danych osobowych w procesie oceny zdolności kredytowej

3.1. Regulacje prawne

Obecnie brak w Polsce regulacji prawnych, które wprost odnosiłyby się do wykorzystywania systemów sztucznej inteligencji, w szczególności w procesach przetwarzania danych osobowych. Nie ulega jednak wątpliwości, że faktycznie systemy

zdefiniowane w punkcie 1. niniejszego opracowania są stosowane np. do oceny zdolności kredytowej. Jak zostało już wyżej wskazane, w takich sytuacjach dochodzi do przetwarzania danych osobowych. Czy więc tego typu procesy podlegają jakimkolwiek wymaganiom prawnym i są w jakiś sposób kontrolowane? Oczywiście odpowiedź na to pytanie brzmi: tak.

Choć ogólne rozporządzenie o ochronie danych osobowych (RODO)[5] jest neutralne technologicznie i nie odnosi się wprost do konkretnych rodzajów rozwiązań technicznych, zawarte w nim postanowienia bez wątpienia można odnieść do systemów sztucznej inteligencji przetwarzających dane osobowe.

Po pierwsze należy wskazać na art. 4 pkt 4 RODO, zawierający definicję profilowania. Jest to zautomatyzowane przetwarzanie danych osobowych służące do oceny niektórych czynników osobowych osoby fizycznej (np. do analizy lub prognozy aspektów dotyczących efektów pracy tej osoby fizycznej, jej sytuacji ekonomicznej, zdrowia, osobistych preferencji, zainteresowań, wiarygodności, zachowania, lokalizacji lub przemieszczania się). Zestawiając tę definicję z definicją systemu sztucznej inteligencji dochodzimy do wniosku, że algorytmy sztucznej inteligencji analizujące dane osobowe mogą dokonywać profilowania. Profilowanie może z kolei skutkować podjęciem decyzji wyłącznie przez maszynę, jak również może być jedynie elementem szerszego procesu, w którym dodatkowo bierze udział człowiek.

Ocena zdolności kredytowej ewidentnie obejmuje profilowanie. Na podstawie dostarczonych informacji o kliencie banku lub innej instytucji dokonywana jest analiza dotycząca zdolności danej osoby do spłaty zobowiązania. Trudno sobie wyobrazić, aby aktualnie takie procesy były dokonywane wyłącznie w sposób manualny. Z uwagi na ilość danych poddawanych analizie oraz liczbę klientów danej instytucji, procesy takie są zautomatyzowane.

Poza profilowaniem RODO mówi jeszcze w art. 22 ust. 1 o zautomatyzowanym podejmowaniu decyzji, a więc takim, w którym człowiek nie uczestniczy. Wskazany przepis ma zastosowanie tylko do takich zautomatyzowanych decyzji, które wywołują wobec danej osoby skutki prawne lub w podobny sposób istotnie na nią wpływają. Co istotne, zgodnie

z RODO, co do zasady podmioty danych mają prawo do tego, by nie podlegać zautomatyzowanym decyzjom. Powyższe nie ma zastosowania w przypadkach, gdy decyzja:

- jest niezbędna do zawarcia lub wykonania umowy między osobą, której dane dotyczą, a administratorem;
- jest dozwolona prawem Unii lub prawem państwa członkowskiego, któremu podlega administrator i które przewiduje właściwe środki ochrony praw, wolności i prawnie uzasadnionych interesów osoby, której dane dotyczą; lub
- opiera się na wyraźnej zgodzie osoby, której dane dotyczą.

Mając na uwadze powyższe regulacje należy się zastanowić, czy faktycznie mogą one znaleźć zastosowanie do procesu oceny zdolności kredytowej. Jak już zostało wskazane, jeśli człowiek nie uczestniczy w procesie analizy, lecz ocena, że klient będzie mógł lub nie będzie w stanie spłacić zaciągniętego kredytu, dokonywana jest na podstawie wyniku dostarczonego przez system sztucznej inteligencji, pierwsza z przesłanek zastosowania art. 22 RODO jest spełniona. Drugim elementem koniecznym do zastosowania przywołanego przepisu jest istotny wpływ decyzji na daną osobę. Nie ulega wątpliwości, że decyzja dotycząca przyznania danej osobie kredytu bądź odmowy udzielenia finansowania, ma bardzo duże znaczenie dla podmiotu danych. Zdecydowanie więc można przyjąć, że analiza zdolności kredytowej i płynące z niej wnioski mieszczą się w definicji zautomatyzowanej decyzji, o której mowa w art. 22 ust. RODO.

3.2. Podstawa prawna przetwarzania

Skoro ocena zdolności kredytowej może podpadać pod definicję wskazaną w art. 22 ust. 1 RODO, istotna jest odpowiedź na pytanie, na jakiej podstawie prawnej dozwolone jest przetwarzanie danych osobowych w celu zautomatyzowanego podjęcia decyzji w wyniku powyższej analizy. Przepisy prawa wprost nakazują bankom lub innym instytucjom dokonanie analizy kredytowej. Zgodnie z art. 70 ust. 1 Prawa bankowego[6]:

"Bank uzależnia przyznanie kredytu od zdolności kredytowej kredytobiorcy." Z kolei art. 9 ust. 1 ustawy o kredycie konsumenckim[7] przewiduje: "Kredytodawca przed zawarciem umowy o kredyt konsumencki jest zobowiązany do dokonania oceny zdolności kredytowej konsumenta." W obu przepisach mowa jest jednak o samej analizie zdolności kredytowej, a nie o zautomatyzowanym podejmowaniu decyzji. Teoretycznie więc, jeśli nie zachodziłby jeden z wyjątków wskazanych w art. 22 ust. 2 RODO, podmiot danych mógłby sprzeciwić się podejmowaniu wobec niego zautomatyzowanej decyzji w omawianym zakresie. Należy więc przeanalizować, czy któraś ze wskazanych w przywołanym przepisie przesłanek ma jednak zastosowanie.

Jedną z opcji dopuszczających zautomatyzowane podejmowanie decyzji jest niezbędność do zawarcia lub wykonania umowy między osobą, której dane dotyczą, a administratorem. Bez wątpienia bez dokonania oceny zdolności kredytowej bank lub inna instytucja nie udzieli kredytu lub pożyczki. Wydawałoby się więc na pierwszy rzut oka, że ta podstawa prawna mogłaby znaleźć zastosowanie. Z RODO wynika jednak, że to zautomatyzowana decyzja ma być niezbędna do zawarcia lub wykonania umowy. Czy tak w istocie jest? Czy udział człowieka w procesie analizy kredytowej nie jest możliwy? Oczywiście, odpowiedź brzmi: nie. Decyzja stanowiąca wynik analizy zdolności kredytowej jak najbardziej może być podjęta w procesie mieszanym, tj. takim, gdzie wnioski dostarczone przez algorytm są następnie poddawane analizie człowieka. Omawiany wyjątek z art. 22 ust. 2 RODO nie może mieć więc zastosowania dla zautomatyzowanego podejmowania decyzji w procesie oceny zdolności kredytowej.

Innym rozwiązaniem mogłaby być wyraźna zgoda osoby, której dane dotyczą. Jak wynika z art. 4 pkt 11 RODO zgoda to dobrowolne, konkretne, świadome i jednoznaczne okazanie woli. Zgoda nie może być wymuszona. Podmiot danych powinien mieć swobodę w zakresie udzielenia zgody, odmowy jej udzielenia bądź jej wycofania, bez jakichkolwiek negatywnych konsekwencji. Oceniając, czy zgodę wyrażono dobrowolnie, w jak największym stopniu uwzględnia się, czy między innymi od zgody na przetwarzanie danych nie jest uzależnione wykonanie umowy, w tym świadczenie usługi, jeśli przetwarzanie danych

osobowych nie jest niezbędne do wykonania tej umowy (art. 7 ust. 4 RODO). Analiza wskazanych wymagań dotyczących wyrażenia zgody prowadzi do wniosku, że zautomatyzowane podejmowanie decyzji w procesie analizy zdolności kredytowej nie może opierać się na wskazanej podstawie prawnej. Podmiot danych nie ma bowiem w rzeczywistości swobody w zakresie udzielenia zgody bądź jej odmowy czy wycofania swojego oświadczenia. Wymaganie udzielenia zgody klienta na zautomatyzowaną ocenę zdolności kredytowej nie zapewniałoby więc dobrowolności.

Ostatnią opcją, jaka ewentualnie pozostaje dla stosowania zautomatyzowanego podejmowania decyzji w procesie oceny zdolności kredytowej, jest dopuszczenie takiej możliwości przez prawo. Tak w rzeczywistości jest, tj. art. 105a Prawa bankowego przewiduje w ust. 1a, że banki, inne instytucje ustawowo upoważnione do udzielania kredytów, instytucje pożyczkowe oraz podmioty, o których mowa w art. 59d ustawy z dnia 12 maja 2011 r. o kredycie konsumenckim, a także instytucje utworzone na podstawie art. 105 ust. 4, mogą w celu oceny zdolności kredytowej i analizy ryzyka kredytowego podejmować decyzje, opierając się wyłącznie na zautomatyzowanym przetwarzaniu, w tym profilowaniu, danych osobowych – również stanowiących tajemnicę bankową – pod warunkiem zapewnienia osobie, której dotyczy decyzja podejmowana w sposób zautomatyzowany, prawa do otrzymania stosownych wyjaśnień co do podstaw podjętej decyzji, do uzyskania interwencji ludzkiej w celu podjęcia ponownej decyzji oraz do wyrażenia własnego stanowiska. Jednocześnie w art. 105a ust. 1b podkreślono, że powyższe decyzje mogą być podejmowane wyłącznie w oparciu o dane niezbędne z uwagi na cel i rodzaj kredytu.

Jak wynika z powyższych regulacji, przepisy prawa uprawniają więc wymienione w nich podmioty do zautomatyzowanego podejmowania decyzji w procesie oceny zdolności kredytowej. Takie czynności podlegają jednak określonym ograniczeniom i wymagają spełnienia dodatkowych obowiązków w porównaniu z analizami kończącymi się decyzjami podejmowanymi z udziałem człowieka. Jak wynika wprost z art. 105a ust. 1 b Prawa bankowego, zautomatyzowane decyzje nie mogą być podejmowane na podstawie analizy dowolnych danych, w oderwaniu od rodzaju kredytu, o który dana osoba wnioskuje. Wręcz

przeciwnie, należy ograniczyć zakres danych poddawanych analizie do niezbędnego minimum.

3.3. Obowiązek informacyjny

Jeśli decyzja o udzieleniu kredytu nie jest podejmowana wyłącznie na podstawie wniosku dostarczonego przez system sztucznej inteligencji, lecz mamy w procesie czynnik ludzki, na instytucji dokonującej oceny zdolności kredytowej ciążyą takie same obowiązki, jak w każdym innym procesie przetwarzania danych osobowych. Administrator przede wszystkim ma obowiązek poinformować daną osobę o przetwarzaniu jej danych osobowych, zgodnie z art. 13 oraz ewentualnie art. 14 RODO.

Zdecydowanie więcej obowiązków ciąży na banku czy innej instytucji, w której ocena zdolności kredytowej obejmuje zautomatyzowane podejmowanie decyzji. Obowiązek informacyjny przekazywany klientowi w związku z przetwarzaniem jego danych osobowych w powyższym procesie powinien obejmować w szczególności informacje o:

- a) samym fakcie zautomatyzowanego podejmowania decyzji i zasadach ich podejmowania (chodzi o istotne informacje), a także o znaczeniu i przewidywanych konsekwencjach takiego przetwarzania dla osoby, której dane dotyczą;
- b) prawie do otrzymania stosownych wyjaśnień co do podstaw podjętej decyzji;
- c) prawie do uzyskania interwencji ludzkiej w celu podjęcia ponownej decyzji oraz do wyrażenia własnego stanowiska.

Można więc wyróżnić uprzednie obowiązki informacyjne danego administratora podejmującego decyzje wyłącznie na podstawie zautomatyzowanego przetwarzania (tj. obowiązki, które musi zrealizować przed przystąpieniem do przetwarzania), jak i obowiązki następcze (aktualizujące się po podjęciu decyzji, gdy dana osoba wnosi o wyjaśnienia i interwencję ludzką). W praktyce najwięcej wątpliwości wśród administratorów budzi kwestia szczegółowości, z jaką dany podmiot ma opisać osobie fizycznej zasady

podejmowania zautomatyzowanych decyzji, a następnie wyjaśnić podstawy podjętej, konkretnej decyzji. Główna obawa dotyczy możliwości ujawnienia tajemnicy przedsiębiorstwa i skopiowania przez konkurencję stosowanych przez danego administratora rozwiązań technicznych.

Odnosząc się do powyższego należy podkreślić, że przepisy RODO, Prawa bankowego, ani żadne inne regulacje prawne nie wymagają przedstawiania szczegółów technicznych dotyczących systemów sztucznej inteligencji wykorzystywanych do oceny zdolności kredytowej. Co więcej, przedstawianie takich szczegółów mogłoby spowodować, że podmioty danych niewiele zrozumiałyby z przekazywanej informacji. Jak natomiast podkreśla się, np. w wytycznych brytyjskiego organu nadzorczego (ICO)[8], administrator powinien w sposób transparentny przekazać informacje dotyczące podejmowania decyzji w sposób zautomatyzowany (z wykorzystaniem systemów sztucznej inteligencji). Oznacza to, że należy w sposób jasny, prosty, otwarty i uczciwy informować o tym, w jaki sposób dane są przetwarzane. W celu ustalenia zakresu i sposobu przekazania informacji osobom, których dane dotyczą, należy uwzględnić kontekst przetwarzania (w szczególności, jakich osób dotyczy przetwarzanie, czy występują jakieś szczególne okoliczności w tym zakresie, jak również jakiego rodzaju wpływ ma podejmowana decyzja na daną osobę itp.). Nie jest absolutnie wymagane zdradzanie przez przedsiębiorcę swojego know – how. Nie ma niestety jednego uniwersalnego sposobu informowania o zasadach podejmowania zautomatyzowanych decyzji. Administrator powinien dobrać metodę i zakres informowania adekwatny do konkretnych okoliczności. Przykłady i różne możliwe podejścia do realizacji omawianych obowiązków zostały dość szeroko przedstawione we wspomnianych wyżej wytycznych ICO, warto więc w tym zakresie po nie sięgnąć.

3.4. Ocena ryzyka

Poza koniecznością spełnienia obowiązków informacyjnych, które w przypadku zautomatyzowanego podejmowania decyzji są bardziej złożone niż w innego rodzaju procesach przetwarzania danych osobowych, istotnym obowiązkiem jest dokonanie oceny

ryzyka oraz oceny skutków przetwarzania dla ochrony danych osobowych (DPIA). W tym miejscu należy przypomnieć, że analiza ryzyka to proces wymagany dla każdego procesu przetwarzania danych osobowych. Wynika to z art. 24 ust. 1 RODO oraz art. 32 ust. 1 RODO, zgodnie z którymi odpowiednie środki techniczne i organizacyjne są wdrażane przez administratora z uwzględnieniem m. in. ryzyka naruszenia praw i wolności osoby fizycznej. Środki te mają bowiem zapewnić stopień bezpieczeństwa danych osobowych odpowiadający ustalonemu ryzyku.

W przypadku stwierdzenia na podstawie powyższej analizy, że dany rodzaj przetwarzania – w szczególności z użyciem nowych technologii – ze względu na swój charakter, zakres, kontekst i cele z dużym prawdopodobieństwem może powodować wysokie ryzyko naruszenia praw lub wolności osób fizycznych, administrator przed rozpoczęciem przetwarzania dokonuje oceny skutków planowanych operacji przetwarzania dla ochrony danych osobowych (DPIA). Jest to więc dodatkowy obowiązek, które materializuje się w przypadku ustalenia, że dane procesy przetwarzania danych osobowych mogą powodować wysokie ryzyko. Artykuł 35 ust. 3 RODO wymienia przykłady takich rodzajów przetwarzania danych osobowych, wymagających przeprowadzenia DPIA. Wśród nich znajduje się dokonywanie systematycznej, kompleksowej oceny czynników osobowych odnoszących się do osób fizycznych, która opiera się na zautomatyzowanym przetwarzaniu, w tym profilowaniu, i jest podstawą decyzji wywołujących skutki prawne wobec osoby fizycznej lub w podobny sposób znacząco wpływających na osobę fizyczną. Analiza zdolności kredytowej i podejmowanie zautomatyzowanych decyzji w tym zakresie mieści się w powyższej definicji, wobec czego bez wątpienia tego typu procesy przetwarzania danych osobowych wymagają dokonania DPIA. Dodatkowo organy nadzorcze w poszczególnych państwach członkowskich mogą publikować listy procesów wymagających dokonania oceny skutków dla ochrony danych osobowych. Prezes Urzędu Ochrony Danych Osobowych wydał komunikat zawierający listę takich procesów[9]. Znalazła się w nim m. in. ocena zdolności kredytowej, przy użyciu algorytmów sztucznej inteligencji, objęta

obowiązkiem zachowania tajemnicy i żądanie ujawnienia danych niemających bezpośredniego związku z oceną zdolności kredytowej.

Ocena skutków planowanych operacji przetwarzania dla ochrony danych osobowych powinna mieć miejsce przede wszystkim przed rozpoczęciem planowanych procesów przetwarzania. W toku realizowania określonych czynności przetwarzania, w razie potrzeby, a przynajmniej gdy zmienia się ryzyko wynikające z operacji przetwarzania, administrator dokonuje przeglądu, by stwierdzić, czy przetwarzanie odbywa się zgodnie z oceną skutków dla ochrony danych. Należy pamiętać, że jeżeli DPIA wskaże, że przetwarzanie powodowałoby wysokie ryzyko, a administrator nie zastosuje środków w celu zminimalizowania tego ryzyka, to przed rozpoczęciem przetwarzania administrator konsultuje się z organem nadzorczym.

4. Ocena zdolności kredytowej w projekcie rozporządzenia dot. sztucznej inteligencji

Warto wspomnieć, że zgodnie z projektem rozporządzenia Parlamentu i Rady w sprawie sztucznej inteligencji, systemy sztucznej inteligencji wykorzystywane do oceny zdolności kredytowej są kwalifikowane jako systemy wysokiego ryzyka. Oznacza to, że ich wykorzystywanie najprawdopodobniej (o ile regulacja zostanie przyjęta w aktualnie opublikowanym kształcie) będzie wiązało się z koniecznością spełnienia szeregu wymagań. Zgodnie z ww. projektem z korzystaniem z systemów wysokiego ryzyka będzie wiązać się m.in.:

- nadzór wyznaczonego organu państwowego,
- obowiązki przejrzystości - odpowiednie informowanie użytkowników systemów sztucznej inteligencji, ale z zastrzeżeniem, że nie będzie to naruszać praw własności intelektualnej,
- wprowadzenie odpowiednich metod zarządzania danymi,
- prowadzenie odpowiedniej dokumentacji technicznej, dzienników logów,

- zapewnienie nadzoru ludzkiego,
- zapewnienie rzetelności, odporności, cyberbezpieczeństwa,
- wdrożenie systemu zarządzania jakością,
- obowiązki rejestracyjne.

5. Podsumowanie.

Wykorzystywanie systemów sztucznej inteligencji, jak również samo pojęcie sztucznej inteligencji, nie są uregulowane wprost w powszechnie obowiązujących przepisach prawa. Powszechnie przyjmuje się jednak szerokie rozumienie powyższego pojęcia, jako mechanizmów i technik analizujących (opartych na wykorzystaniu algorytmów), interpretujących zdarzenia i informacje (dane), w celu automatyzacji podejmowania decyzji i wykonywania określonych działań.

Systemy sztucznej inteligencji są wykorzystywane m. in. w sektorze finansowym dla oceny zdolności kredytowej. Omawiany proces wymaga analizy wielu informacji stanowiących dane osobowe, a w jej wyniku są wytwarzane nowe dane osobowe. Konieczne jest więc w tym zakresie przestrzeganie wymagań RODO i innych powszechnie obowiązujących przepisów z zakresu ochrony danych osobowych. W szczególności zastosowanie w omawianym zakresie będą miały wymagania związane ze zautomatyzowanym podejmowaniem decyzji w rozumieniu art. 22 ust. 1 RODO. Przepisy Prawa bankowego uprawniają banki i inne instytucje do zautomatyzowanego dokonywania oceny zdolności kredytowej. Wiąże się to z koniecznością informowania osób wnioskujących o kredyt o zasadach zautomatyzowanego podejmowania decyzji, prawie do uzyskania wyjaśnienia podstaw podjętej decyzji, jak również o prawie do uzyskania interwencji ludzkiej i wyrażenia własnego stanowiska przez osobę, której dane dotyczą. Informacje powinny być przekazywane podmiotom danych w sposób prosty, jasny, bez zbędnych szczegółów

technicznych. W szczególności nie jest wymagane podawanie przez administratorów informacji stanowiących tajemnicę przedsiębiorstwa. Zautomatyzowana ocena zdolności kredytowej, jako proces powodujący wysokie ryzyko dla praw i osoby fizycznej, wymaga także przeprowadzenia DPIA.

Wykorzystywanie systemów sztucznej inteligencji, w tym w szczególności dokonywanie oceny zdolności kredytowej, prawdopodobnie zostanie wyraźnie uregulowane w rozporządzeniu dotyczącym sztucznej inteligencji, którego projekt został opublikowany w kwietniu br. Systemy stosowane do analizy zdolności kredytowej będą uważane za systemy wysokiego ryzyka, co spowoduje konieczność wypełniania obowiązków nie tylko na podstawie RODO, ale również wymagań wynikających z planowanego aktu prawnego.

[1] Recommendation of the Council on Artificial Intelligence, https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449?_ga=2.67738585.97980544.1621678804-142485567.1621678804, dostęp: 22.05.2021 r.

[2] Uchwała nr 196 Rady Ministrów z dnia 28 grudnia 2020 r. w sprawie ustanowienia „Polityki dla rozwoju sztucznej inteligencji w Polsce od roku 2020”, Monitor Polski, poz. 23.

[3] Proposal for a Regulation of the European Parliament and of the Council, Laying down harmonised rules on artificial intelligence (artificial intelligence act) and amending certain union legislative acts, COM (2021) 206 final, 21.04.2021

[4] Tłum. własne.

[5] Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2016/679 z dnia 27 kwietnia 2016 r. w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych i w

sprawie swobodnego przepływu takich danych oraz uchylenia dyrektywy 95/46/WE (ogólne rozporządzenie o ochronie danych), L119/1.

[6] U S T AWA z dnia 29 sierpnia 1997 r. Prawo bankowe, Dz. U. 1997 Nr 140 poz. 939 ze zm.

[7] U S T AWA dnia 12 maja 2011 r. o kredycie konsumenckim, Dz. U. 2011 Nr 126 poz. 715 ze zm.

[8] Explaining decisions made with AI, Information Commissioner's Office, 20.05.2020, <https://ico.org.uk/for-organisations/guide-to-data-protection/key-data-protection-themes/explaining-decisions-made-with-ai/>

[9] KOMUNIKAT PREZESA URZĘDU OCHRONY DANYCH OSOBOWYCH z dnia 17 czerwca 2019 r., Monitor Polski 2019, poz. 666

Konstruowanie humanocentrycznej SI z wykorzystaniem instrumentów prawnych ochrony danych osobowych

Autor: dr Dominik Lubasz

Lubasz i Wspólnicy – Kancelaria Radców Prawnych sp. k.

1. Wprowadzenie

Gwałtowny rozwój systemów sztucznej inteligencji i ich potencjalnie bardzo szerokie zastosowanie zrodziły obawy związane z niepożądanymi efektami, zwłaszcza negatywnym wpływem na prawa i wolności człowieka. Obawy te legły u podstaw kształtowania się kierunków poszukiwania rozwiązań mitygujących ryzyka związane z jej tworzeniem i wykorzystywaniem. Jednym z obszarów zainteresowania związanego z oczywistą zależnością pomiędzy zarówno rozwojem, jak i aplikacją rozwiązań SI, a dostępnością, jakością i ilością danych, jest obszar danych osobowych. To na podstawie danych, w tym w szczególności danych osobowych, systemy samouczące realizują swoje pierwotne przeznaczenie, a zatem znalezienie związków pomiędzy informacjami pozwalające na wyciąganie określonych wniosków, podejmowanie decyzji, a w konsekwencji, wywieranie wpływu na otoczenie w obszarze przewidzianym przez twórcę. Oczywiście prawidłowość działania takich mechanizmów analizujących dane, a zatem i potencjalne skutki niepożądane np. arbitralne lub dyskryminujące efekty zautomatyzowanych decyzji podejmowanych z ich wykorzystaniem, zależą również od poprawności samych algorytmów. W tym kontekście bardzo istotne jest uwzględnienie deficytów zaufania do technologii.

2. Dotychczasowe kluczowe działania zmierzające do określenia podejścia regulacyjnego sztucznej inteligencji w Unii Europejskiej

Konieczność kompleksowej analizy związanej z rozwojem sztucznej inteligencji i jej implikacji dla praw podstawowych, wskazana została m.in. w kolejnych opracowaniach organów unijnych, tj. komunikatach Komisji Europejskiej z dnia 25.04.2018 r. „Sztuczna inteligencja dla Europy”[2] oraz w „Skoordynowanym planie w sprawie sztucznej inteligencji” z 7.12.2018 r.[3], a także w „Deklaracji współpracy w sprawie SI” podpisanej 10.4.2018 r.[4] oraz Białej Księdze w sprawie sztucznej inteligencji z 19.2.2020 r.[5], wytycznych Grupy Ekspertów Wysokiego Szczebla ds. Sztucznej Inteligencji (High-Level Expert Group on AI (HLEG))[6], jak również rezolucjach Parlamentu Europejskiego m.in. z dnia 20 października 2020 r. w sprawie praw własności intelektualnej w dziedzinie rozwoju technologii sztucznej inteligencji[7], z dnia 20.10.2020 r. z zaleceniami dla Komisji w sprawie systemu odpowiedzialności cywilnej za sztuczną inteligencję[8] oraz z dnia 20.10.2020 r. zawierające zalecenia dla Komisji w sprawie ram aspektów etycznych sztucznej inteligencji, robotyki i powiązanych z nimi technologii[9].

Wspólnym zamierzeniem ww. opracowań było poszukiwanie kierunków i wymagań jakie powinny spełniać ramy prawne dla rozwoju SI uwzględniając mechanizmy wyważania korzyści związanych z rozwojem technologicznym i aspektów etycznymi skoncentrowanymi na pojęciu humanocentrycznej sztucznej inteligencji, przede wszystkim przez nadanie jej cechy „godnej zaufania” sztucznej inteligencji. Podejście to potwierdziła znalazło najszersze potwierdzenie w opracowanych przez powołaną przez Komisję Europejską Grupę Ekspertów Wysokiego Szczebla ds. Sztucznej Inteligencji Wytycznych w zakresie etyki dotyczące godnej zaufania sztucznej inteligencji[10], a także we wspomnianej Białej księdze z dnia 19.02.2020 r. w sprawie sztucznej inteligencji, w których sformułowane zostały wymogi dla godnej zaufania sztucznej inteligencji. Skupiają się one wokół wyposażenia człowieka poddanego wpływowi procesów wykorzystujących SI w gwarancje autonomii i kontroli, a także ochrony. W obszarze prywatności i ochrony danych osobowych wskazano równocześnie na

zagadnienia wymagające oceny *fit for purpose* obecnej regulacji dotyczącej ochrony danych osobowych i praw konsumenta. Wśród kluczowych punktów problemowych wymieniono:

- 1) sprzeczność między innowacjami w obszarze big data a ochroną danych;
- 2) społeczne i etyczne implikacje big data (takie jak profilowanie i przejrzystość procesu decyzyjnego, stronniczość w danych itp.);
- 3) zapewnienie bezpiecznego i godnego zaufania udostępniania danych osobowych;
- 4) ograniczenia techniczne, w szczególności granice anonimizacji i pseudonimizacji;
- 5) pochodzenie danych z różnych źródeł, w tym niezweryfikowanych czy niezaufanych;
- 6) konieczności zapewnienia globalnego poziomu ochrony danych oraz weryfikacja istnienia łatwo egzekwowalnych instrumentów ochrony;
- 7) podejście oparte na ryzyku wyznaczające zakres obowiązków administratorów danych.

3. Projekt rozporządzenia w sztucznej inteligencji

Prace analityczne na płaszczyźnie unijnej w obszarze regulacji sztucznej inteligencji doprowadziły do przedstawienia przez Komisję Europejską w dniu 21.04.2021 roku projektu rozporządzenia ustanawiającego zharmonizowane przepisy dotyczące sztucznej inteligencji (akt w sprawie sztucznej inteligencji) i zmieniającego niektóre akty ustawodawcze Unii[11]. Rozporządzenie to ma ustanawiać zharmonizowane zasady dotyczące sztucznej inteligencji, przewidując przepisy mające zastosowanie do projektowania, opracowywania i stosowania niektórych systemów SI wysokiego ryzyka oraz ograniczeń dotyczących niektórych zastosowań systemów zdalnej identyfikacji biometrycznej. Wdrożone mają zostać rozwiązania prawne w zakresie niedyskryminacji wprowadzając wymogi, których celem jest

zminimalizowanie ryzyka dyskryminacji algorytmicznej, w szczególności w odniesieniu do projektowania i jakości zbiorów danych wykorzystywanych do opracowywania systemów SI uzupełnionych o obowiązki w zakresie testowania, zarządzania ryzykiem, dokumentacji i nadzoru ludzkiego w całym cyklu życia systemów SI.

Rozporządzenie w sprawie sztucznej inteligencji, pomimo szerokiej definicji systemów SI, obejmuje tylko niektóre aspekty związane z funkcjonowaniem systemów SI i tylko niektóre systemy. Wynika to z faktu, że Komisja Europejska założyła, co podkreślano we wcześniejszych opracowaniach, że w obszarze ochrony konsumentów czy danych osobowych Unia Europejska posiada silne i zrównoważone ramy regulacyjne, które mogą wyznaczać globalny standard podejścia do tej technologii. W związku z tym, zgodnie z pkt 1.2 uzasadnienia, rozporządzenie w sprawie SI ma jedynie uzupełniać ogólne rozporządzenie o ochronie danych i dyrektywę 2016/680 o zestaw zharmonizowanych przepisów mających zastosowanie do projektowania, opracowywania i wykorzystywania niektórych systemów AI wysokiego ryzyka oraz ograniczeń dotyczących niektórych zastosowań systemów zdalnej identyfikacji biometrycznej. Z tych względów RODO ma zatem pozostać główną podstawą oceny zgodności systemów SI w obszarze danych osobowych. Cele tego aktu prawnego polegające na zapewnieniu wysokiego poziomu ochrony praw osób fizycznych, a także mechanizmy prawne oparte na podejściu neutralnym technologicznie, zgodnie z którym realizacja wymogów opiera się na analizie ryzyka przetwarzania z punktu widzenia praw i wolności osób, których dane dotyczą, zapewnienia zgodności z zasadami przetwarzania, są bowiem zgodne z trendami rozwoju technologii oraz założeniami ram prawnych godnej zaufania, humanocentrycznej sztucznej inteligencji.

4. Ogólne rozporządzenie o ochronie danych osobowych jako element ram prawnych dla konstruowania humanocentrycznej sztucznej inteligencji

Przedstawiony kierunek regulacji SI wykorzystuje zbieżność celów tworzenia ram prawnych dla SI z podstawowym celem regulacyjnym rozporządzenia 2016/679 tj.

zapewnienie wysokiego poziomu ochrony praw osób fizycznych oraz ułatwienie swobodnego przepływu danych osobowych na jednolitym rynku cyfrowym. Cel ten ma być osiągnięty m.in. przez metodę regulacji, w szczególności neutralność technologiczną, podejście oparte na ryzyku, tj. uzależnienie wyboru środków technicznych i organizacyjnych mających za zadanie zapewnić zgodność z rozporządzeniem na analizie ryzyka (risk-based approach) oraz przez wprowadzenie rozwiązań opartych na zasadach i klauzulach generalnych, w miejsce kazuistycznych modeli legislacyjnych.

Ogólne ramy systemu ochronnego danych osobowych precyzowane regulacją szczegółową nakładającą na administratorów względny zakaz zautomatyzowanego podejmowania decyzji w (art. 22 RODO), obowiązek uwzględniania ochrony danych w fazie projektowania (*data protection by design* – art. 25 ust. 1 RODO) oraz domyślnej ochrony danych (*data protection by default* – art. 25 ust. 2 RODO), a także konieczność przeprowadzania oceny skutków przetwarzania dla ochrony danych (*data protection impact assessment* – art. 35 RODO). To przede wszystkim te instrumenty legislacyjne z uwzględnieniem zasad przetwarzania danych osobowych (art. 5) będą kluczowe przy ocenie skuteczności obowiązujących już ram prawnych dla rozwoju sztucznej inteligencji.

Zwłaszcza ciekawa regulacyjnie z perspektywy opracowania systemów SI jest koncepcja *data protection by design*, dla której pierwowzorem jest *privacy by design* stworzona przez dr *Ann Cavoukian*. Koncepcja *privacy by design* opiera się ona na siedmiu filarach:

- 1) proaktywnym podejściu, a nie reaktywnym i zaradczym, nie naprawczym;
- 2) prywatności jako ustawieniu domyślnym;
- 3) prywatności włączonej w projekt;
- 4) pełnej funkcjonalności rozumianej jako osiągnięcie sumy dodatniej, nie sumy zerowej;
- 5) ochronie prywatności od początku do końca cyklu życia informacji;

- 6) transparentności i przejrzystości;
- 7) poszanowaniu dla prywatności użytkowników[12].

Warto przy tym zwrócić uwagę, że w kontekście sztucznej inteligencji A. Cavoukian zmodyfikowała pierwotną koncepcję wskazując na potrzebę etycznego konstruowania narzędzi wykorzystujących mechanizmy sztucznej inteligencji (*AI Ethics by design*). Elementami podstawowymi tej koncepcji stały się: przejrzystość i rozliczalność algorytmów; stosowanie zasad etycznych przy przetwarzaniu danych osobowych; zapewnienie nadzoru i odpowiedzialności za działanie algorytmów; poszanowanie prywatności jako podstawowego prawa człowieka; ochrona danych jako ustawienie domyślne; proaktywna identyfikacja zagrożeń dla bezpieczeństwa, a tym samym minimalizacja zagrożeń; solidna dokumentacja ułatwiająca etyczne projektowanie i symetrię danych[13].

Koncepcja Ann Cavoukian zaadaptowana została w ogólnym rozporządzeniu o ochronie danych i służy jego celom ochronnym jakimi jest zapewnienie ochrony praw i wolności podmiotów danych w związku z przetwarzaniem ich danych osobowych, przy uwzględnieniu jednak zasad ochrony danych uregulowanych w art. 5 RODO zwłaszcza rzetelności i przejrzystości, minimalizacji danych oraz mechanizmu *risk-based approach*. Istotą *data protection by design* w rozumieniu art. 25 RODO jest obowiązek administratora uwzględnienia ochrony danych osobowych już w fazie projektowania określonego rozwiązania, usługi czy systemu sztucznej inteligencji. Ma to zagwarantować m.in., że ochrona danych osobowych stanie się immanentnym elementem każdego projektu już na etapie tworzenia. Takie rozwiązania będzie zatem wpisane w każdy z etapów tworzenia technologii i w każdym z etapów od szkolenia począwszy a na wdrożeniu skończywszy skupione będzie na ocenie wpływu na podmioty danych w świetle zasad RODO.

W omawianym instrumencie nie chodzi zatem jedynie o aspekt zgodności, lecz również, a może przede wszystkim, o rozwinięcie określonej kultury organizacyjnej pracy nad nowymi rozwiązaniami, by tę zgodność zapewnić. Wymóg uwzględniania ochrony danych w fazie projektowania, pozwala na zapewnienie zgodności nie tylko z art. 25 RODO, ale także

z pozostałymi wymogami w oparciu o mechanizm analizy ryzyka dla praw i wolności podmiotów danych, który implikuje konieczność rozpoczęcia analizy od weryfikacji zasad przetwarzania, legalizacji, możliwości realizacji uprawnień podmiotów danych, a skończywszy na bezpieczeństwie przetwarzania.

Zwłaszcza ten ostatni element jest charakterystyczny dla regulacji RODO, jako że nakłada obowiązek zapewnienia bezpieczeństwa o charakterze kontekstowym, zależnym o przeprowadzonej w konkretnym stanie faktycznym analizy ryzyka przetwarzania dla praw i wolności podmiotów danych. Dobór odpowiednich środków zabezpieczenia warunkowany jest ich adekwatnością do ryzyk ocenianych z punktu widzenia osób, których dane mają być przetwarzane. Proaktywna identyfikacja zagrożeń dla bezpieczeństwa, a tym samym minimalizacja zagrożeń, techniczna solidność, są walorami przy projektowaniu rozwiązań opartych na SI. Podejście takie należy przy tym wpisać zarówno w proces tworzenia, jak i wykorzystywania mechanizmów SI.

Z perspektywy konstrukcji mechanizmów opartych na sztucznej inteligencji, natężenie trudności związanych z zapewnieniem zgodności może okazać się różne w obrębie poszczególnych obowiązków, poczynsz od zasad przetwarzania. Nieco inne zagadnienia problematyczne występują przy zapewnieniu zgodności z zasadą legalności, rzetelności, przejrzystości, minimalizacji danych czy zasadą poufności i integralności. Przykładowo z w stosunku do zasady minimalizacji kluczowe będzie takie zaprojektowania systemów SI by zapewnić poprawność i niedyskryminujący charakter przetwarzania lub skutków takich operacji, jak i wykluczyć możliwość oparcia na nich nieobiektywnych lub krzywdzących zautomatyzowanych decyzji oddziałujących na prawa i wolności podmiotów danych. W zakresie zasady przejrzystości powstaje pytanie o odpowiedni poziom transparentności algorytmów i zasad ich działania, by zapewnić osobie, która wchodzi w interakcje z systemem SI, z jednej strony wiedzy co do sposobu i w celu przetwarzania jej danych, którego może się rozsądnie spodziewać, a z drugiej możliwości oceny i zakwestionowania decyzji podejmowanych przez system oddziałujących na sferę jej praw i obowiązków. W tym kontekście znaczenie będzie miał przykładowo wybór metody uczenia. Innego rodzaju

informacje będą musiały być udzielane, gdy system SI oparty jest na głębokim uczeniu, zwłaszcza nienadzorowanym lub jedynie częściowo nadzorowanym, a inne gdy uczenie następuje nie na podstawie metod wykorzystujących symboliczne zasady rozumowania. Podkreślenia przy tym wymaga, że dobór metod, które mają być podstawą działania systemów SI z uwzględnieniem wymogu ich wyjaśnialności jest w konsekwencji jednym z elementów oceny zgodności z zasadą przejrzystości, a trudności technologiczne nie mogą zwalniać z obowiązku zachowania przejrzystości.

Ciężar problemów w zakresie zapewnienia zgodności będzie rozkładał się inaczej w fazie uczenia, testowania, walidacji i inaczej w fazie aplikacji. Jest to szczególnie widoczne w przy realizacji zasady celowości, gdzie każda z faz realizowana jest w innym celu co wpływa równocześnie na zakres danych, ograniczenia retencyjne czy środki bezpieczeństwa. Jest to również zauważalne przy projektowaniu i wykonywaniu obowiązków związanych z realizacją praw podmiotów danych, określonych w art. 12-22 RODO. Zakres poszczególnych praw, możliwość skorzystania z nich przez podmioty danych, zakres obowiązków informacyjnych czy wreszcie przesłanki przełamania względnego zakazu zautomatyzowanego podejmowania decyzji wynikające z art. 22 ust. 2 RODO, będą oceniane w kontekście poszczególnych faz nieco inaczej.

Specyfika technologii wykorzystywanych w związku z tworzeniem systemów SI, powoduje wreszcie, że w niektórych przypadkach obowiązkowe, a w innych wskazane jest wdrożenie szczególnego wymogu przeprowadzenia oceny skutków dla ochrony danych (art. 35 ust. 1, 3 i 4 RODO). Realizacja tego obowiązku polega na ocenie wpływu, jaki projektowane rozwiązanie lub operacja przetwarzania z wykorzystaniem systemu SI może mieć na prawa i wolności osób fizycznych, których dane są lub będą wykorzystywane. Pogłębienie analizy w tym trybie jest jednym z elementów, które mogą być kluczowe dla wykazania spełnienia zasad przetwarzania a w konsekwencji oceny dopuszczalności z perspektywy wymogów RODO. Pozwala równocześnie na podjęcie uprzednich działań zmierzających do ograniczenia negatywnego wpływu przetwarzania na osoby, których dane dotyczą.

5. Podsumowanie

Dyskusja na temat stworzenia ram prawnych dla godnej zaufania humanocentrycznej sztucznej inteligencji weszła obecnie w gorącą fazę, a przedstawiony projekt Komisji Europejskiej należy poczytywać jako zachętę do niej. Jednym z kluczowych elementów tej dyskusji będzie nadal relacja rozporządzenia w sprawie SI do RODO, i możliwości jego skutecznego zastosowania z perspektywy celów ram prawnych dla SI. Już jednak obecnie, po wstępnych analizach przedstawionego projektu rozporządzenia w sprawie SI, pojawiają się głosy o niewystarczającej ochronie podmiotów danych, zwłaszcza w aspektach dotyczących kontroli i przejrzystości, braku zapowiadanej horyzontalności aktu i stworzenia ram prawnych dla wszystkich systemów SI a nie wyłącznie wybranych. Równocześnie przyjęta koncepcja wprowadzania zakazów określonych form wykorzystania SI a priori również może budzić kontrowersje.

Przy konstruowaniu systemów SI, do czasu przyjęcia szczegółowych rozwiązań, z perspektywy ochrony danych osobowych pierwszoplanowe zastosowanie będzie znajdowała regulacja ogólnego rozporządzenia o ochronie danych, zwłaszcza mechanizmy oparte na ryzyku w szczególności wskazane w niniejszym artykule *data protection by design*, czy *data protection impact assessment*, których rozwiązanie kontekstowe o charakterze funkcjonalnym należy wpisać w proces zarówno projektowania, jak i stosowania systemów SI.

Pewne wsparcie o szerszym zakresie, w szczególności umożliwiające ocenę zgodności z wymogami dotyczącymi godnej zaufania SI sformułowanymi w wytycznych HLEG[14], tj. humanocentryczności, nadzoru człowieka, solidności technicznej i bezpieczeństwa, prywatności i zarządzanie danymi, przejrzystości, różnorodności, niedyskryminacji i sprawiedliwości a także dobrobytowi społecznemu i środowiskowemu, i wreszcie rozliczalności, może stanowić również opracowane przez Grupę narzędzie ALTAI - *The Assessment List on Trustworthy Artificial Intelligence*[15].

[1] Autor jest radcą prawnym, doktorem nauk prawnych i współnikiem zarządzającym w Lubasz i Wspólnicy – Kancelarii Radców Prawnych sp.k. Redaktorem i współautorem licznych publikacji z zakresu ochrony danych osobowych oraz prawa nowych technologii, w tym m.in.: Meritum. Ochrona danych osobowych, komentarzy do ogólnego rozporządzenia o ochronie danych, ustawy o prawach konsumenta, o świadczeniu usług drogą elektroniczną, ustawy o ochronie niektórych usług świadczonych drogą elektroniczną opartych lub polegających na dostępie warunkowym, ustawy o ochronie baz danych. ORCID 0000-0001-9716-5802

[2] Komunikat Komisji do Parlamentu Europejskiego, Rady Europejskiej, Rady, Europejskiego Komitetu Ekonomiczno-Społecznego i Komitetu Regionów, Sztuczna inteligencja dla Europy, COM(2018) 237 final.

[3] Komunikat Komisji do Parlamentu Europejskiego, Rady Europejskiej, Rady, Europejskiego Komitetu ekonomiczno-społecznego i komitetu regionów, Skoordynowany plan w sprawie sztucznej inteligencji, COM(2018) 795 final.

[4] Zob. <https://ec.europa.eu/digital-single-market/en/news/eu-member-states-sign-cooperate-artificial-intelligence> (sprawdzono: 9.6.2019 r.).

[5] White Paper On Artificial Intelligence – A European approach to excellence and trust – COM (2020) 65 final.

[6]

https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_PL.pdf (sprawdzono: 9.6.2019 r.).

[7] 2020/2015(INI)

[8] 2020/2014(INL)

[9] 2020/2012(INL)

[10]

https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_PL.pdf (sprawdzono: 9.6.2019 r.)

[11] COM(2021) 206 final.

[12] *Privacy by Design. The 7 Foundational Principles*, dokument dostępny na stronie internetowej: <https://www.ipc.on.ca/wp-content/uploads/Resources/7foundationalprinciples.pdf>, dostęp: 9.06.2021. Zob. także D. Lubasz, K. Witkowska, *RODO. Ogólne rozporządzenie o ochronie danych. Komentarz*, LEX/el., komentarz do art. 25 uwaga 5, zob. też W.R. Wiewiórowski, *Privacy by design jako paradygmat ochrony prywatności* [w:] *Internet. Prawno-informatyczne problemy sieci, portali i e-usług*, red. G. Szpor i W.R. Wiewiórowski, Warszawa 2012, s. 13–29 i wskazaną tam literaturę oraz M. Bienias, *Ochrona danych w fazie projektowania oraz domyślna ochrona danych (privacy by design oraz privacy by default) w ogólnym rozporządzeniu o ochronie danych* [w:] *Ogólne rozporządzenie o ochronie danych. Aktualne problemy prawnej ochrony danych osobowych 2016*, red. G. Sibiga, Warszawa 2016, s. 53–57

[13] A. Cavoukian, *Ethics by design*, dostępne na: www.ryerson.ca/pbdce dostęp: 9.06.2021

[14]

https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_PL.pdf (sprawdzono: 9.6.2019 r.).

[15] <https://futurium.ec.europa.eu/en/european-ai-alliance/pages/altai-assessment-list-trustworthy-artificial-intelligence>, (sprawdzono: 9.6.2019 r.).

Jaki plan dla Sztucznej Inteligencji ma Unia Europejska?

Autorzy: r.pr. Maciej Gawroński, Maciej Grzesiuk

Kancelaria Gawroński & Partners

Korzystanie z technologii Sztucznej Inteligencji staje się światowym standardem i przyszedł chyba najwyższy czas, żeby towarzyszyły temu konkretne uregulowania prawne. Dostrzega to unijny prawodawca, który od paru lat prowadził rozległe badania i konsultacje, celem ustalenia kształtu i kierunku regulacji. Efektem tych prac jest opublikowana 21 kwietnia 2021 r. propozycja Komisji Europejskiej dotycząca Rozporządzenia w sprawie Sztucznej Inteligencji[1].

1. Czym jest Sztuczna Inteligencja?

Pytanie „Czym jest Sztuczna Inteligencja” jest kluczowe do prowadzenia jakichkolwiek rozmów na ten temat – i tylko pozornie proste. Trzeba do niego podchodzić z kilku perspektyw.

W pierwszej kolejności wysoce istotne jest to, jak przedstawiana jest Sztuczna Inteligencja w kulturze popularnej. Można pokusić się wręcz o tezę, że popkultura zdominowała dyskurs o SI. Długie lata upłynęły od momentu, w którym naukowcy zaczęli snuć teorie dotyczące stworzenia „myślących maszyn”, do momentu, w którym takie maszyny zaczęły nieśmiało funkcjonować. W międzyczasie powstało mnóstwo filmów, książek czy gier, które prezentowały swoją interpretację Sztucznej Inteligencji.

Ufamy, że w tym momencie czytelnik ma już przed oczami swoją „ulubioną” popkulturową Sztuczną Inteligencję[2]. Spodziewamy się, że została ona przedstawiona w sposób ludzko-rozrywkowy. Oznacza to po pierwsze, że popkulturowa SI jest stworzona na wzór człowieka: humanoidalny robot, albo przynajmniej mówiący po angielsku superkomputer. Nie tak to wygląda w rzeczywistości i prawdopodobnie jeszcze przez długi czas (a możliwe nigdy) tak nie będzie wyglądało. Być może ktoś jeszcze pamięta androida Sophia, który w 2017 roku został obywatelem Arabii Saudyjskiej – poprzestańmy na tym, że cztery lata później Sophia nie jest oczywiście świadomym uczestnikiem saudyjskiej demokracji[3].

Po drugie, popkulturowa SI co do zasady stanowi (mniejsze lub większe) zagrożenie dla gatunku ludzkiego. Inteligentne maszyny z *Matrix* czy Skynet z *Terminatora* to oczywiście absolutnie klasyczne przykłady takiego podejścia, ale nawet Poe, Sztuczna Inteligencja wspomagająca głównego bohatera serialu *Altered Carbon*, określana jest jako „dzika SI” – ponieważ w świecie *Altered Carbon* Sztuczne Inteligencje dokonały buntu i trzeba je było odgrodzić od reszty sieci. Nie sposób nie zauważyć przekładania się obaw popkulturowych na rzeczywistość. Kiedy Elon Musk (nie bójmy się nazwać go „influencerem” w tej dziedzinie) przedstawia perspektywę ludzkości *niszczonej* przez SI, snuje filmowe wizje apokalipsy. Elon Musk nie skupia się przy tym na aktualnych zagrożeniach wynikających z SI, które nie są zagrożeniami dla przetrwania gatunku ludzkiego a raczej dla jego praw i wolności.

W dalszej kolejności należy się krótko pochylić nad podejściem naukowym. Ponieważ Sztuczna Inteligencja jako dziedzina badań naukowym działa na pograniczu informatyki i kognitywistyki (czerpiąc przy tym m. in. z matematyki i filozofii), definicji naukowych namnożyło się mnóstwo. Z prawnego punktu widzenia rekomendowane jest opieranie się na tych bardziej ogólnych. Najlepsza będzie ta sformułowana przez jednego z ojców nauki o SI, Marviną Minsky’ego: [SI to] *nauka zajmująca się tworzeniem maszyn wykonujących czynności, które wymagałyby inteligencji, gdyby były wykonywane przez ludzi*[4]. Jeśli słowa dotyczące technologii, wypowiedziane ponad pół wieku temu, znajdują zastosowanie również dziś, to mamy do czynienia z dobrą definicją. To podejście wydaje się odpowiadać

temu, co o SI mówi Kate Crawford – główny badacz SI w Microsoft Research: „Sztuczna Inteligencja nie jest ani sztuczna ani inteligentna”[5].

W trzeciej kolejności dochodzimy wreszcie do definicji prawnej, proponowanej w Rozporządzeniu AI. Definicja biegnie dwutorowo:

„system sztucznej inteligencji” oznacza oprogramowanie opracowane przy użyciu co najmniej jednej spośród technik i podejść wymienionych w załączniku I, które może – dla danego zestawu celów określonych przez człowieka – generować wyniki, takie jak treści, przewidywania, zalecenia lub decyzje wpływające na środowiska, z którymi wchodzi w interakcję.

Jak wynika z treści definicji, musi uzupełniać ją *Aneks I*, zawierający listę ogólnie opisanych technik informatycznych. Nie ma w tym miejscu potrzeby jej przytaczania: lista technik została umieszczona w *Aneksie*, żeby mogła ulegać potrzebnym zmianom, a ich poszczególne nazwy niewiele mówią laikom (wydaje się, że na polu identyfikowania konieczna będzie ścisła współpraca prawników i architektów rozwiązań). Oczywiście znalazły się wśród nich rozwiązania typowo kojarzone z SI, jak uczenie maszynowe czy sieci neuronowe.

Definicja w prosty sposób wskazuje więc, że w przypadku systemu SI mówimy o oprogramowaniu, które poprzez wykorzystanie danej techniki lub podejścia generuje wyniki prac. Jeśli natomiast kogoś niepokoi sformułowanie *dla danego zestawu celów określonych przez człowieka* spieszmy wy tłumaczyć, że nawet jeśli SI została zaprogramowana by samej sobie wyznaczać nowe cele, to jest to cel, który został jej wskazany przez człowieka na etapie tworzenia (co ponownie uderza w mit „autonomicznej” SI).

Propozycję definicji zawartą w Rozporządzeniu należy ocenić pozytywnie, jako na tyle konkretną, żeby było jasnym, do jakich systemów się stosuje, a równocześnie na tyle ogólną, by była gotowa na dynamiczny rozwój tej technologii w przyszłości. Zamiast usiłować zgeneralizować zagadnienie, lepiej jest wskazać przykłady technologii, które za SI będziemy

uważać i częściowo oddać definicję rynkowi. Jak wskazują autorzy projektu, definicja została opracowana *w taki sposób, aby w możliwie największym stopniu była neutralna pod względem technologicznym i nie ulegała dezaktualizacji, biorąc pod uwagę szybki rozwój technologiczny i rozwój sytuacji rynkowej w dziedzinie AI*. Ostatecznie wszystko zweryfikuje czas, ale samo podejście wydaje się słuszne.

2. Jak Rozporządzenie dzieli systemy SI?

Skoro już ustaliliśmy, czym jest (oraz częściowo czym nie jest) system Sztucznej Inteligencji z perspektywy prawnej, należy wyjaśnić, jak Rozporządzenie te systemy dzieli. Klasyfikacja SI jest kluczowym elementem Rozporządzenia, ponieważ waży nie tylko na obowiązkach podmiotów wykorzystujących SI, ale również w ogóle na możliwości wprowadzenia SI do obrotu.

Nadmienić trzeba jedynie, że Rozporządzenie nie reguluje militarnego zastosowania Sztucznej Inteligencji.

Jako pierwsze Rozporządzenie wyszczególnia *zakazane użycia Sztucznej Inteligencji*[6], czyli po prostu zakazane systemy SI. W ten sposób zakazane zostaje korzystanie z SI wykorzystującej techniki podprogowe[7] oraz wykorzystującej szczególne podatności danej grupy osób (związane np. z wiekiem), które mogłyby wpłynąć na zachowanie osoby wchodzącej w interakcję z SI tak, że byłoby zagrożone zdrowie fizyczne lub psychiczne tej osoby. Zakazane jest również wykorzystywanie systemów SI służących identyfikacji biometrycznej w czasie rzeczywistym – ale tylko przez służby policyjne (o czym jeszcze w dalszej części opracowania) oraz z szerokimi wyjątkami. W praktyce propozycje zakazów w RAI sformułowane są w najlepszym razie w sposób uznaniowy a w najgorszym fasadowy.

Najciekawszy i wyczekiwany zakaz dotyczy skomplikowanie opisanych systemów SI: używanych przez władze publiczne do ewaluowania lub klasyfikowania osób fizycznych bazując na ich zachowaniach społecznych albo znanych lub przewidywanych cechach

osobistych, prowadzącego do gorszego traktowania tych osób. Powyższy opis (w tekście Rozporządzenia jeszcze bardziej rozbudowany) dotyczy oczywiście niczego innego, jak systemów oceny obywateli, takich jak niesławny chiński system zaufania[8] społecznego, który w przypadku obniżenia „zaufania” do obywatela (np. przez niespłaconą pożyczkę) utrudnia mu dostęp do świadczeń publicznych (np. wyższej edukacji). Systemy tego typu przed nadejściem pandemii COVID-19 były w oczywisty sposób poważnym zagrożeniem dla podstawowych praw i wolności obywatela Unii. Równocześnie jest kwestią czasu, aż takie systemy zaczną być importowane przez autorytarne kraje z całego świata. Komisja Europejska, proponując takie postanowienia, teoretycznie ucina perspektywy importu do krajów UE takich autorytarnych wzorców. Niestety, rozwój legislacji covidowej w UE zmierza wyraźnie w kierunku wdrożenia drastycznej wersji scoringu społecznego.

Główną oś Rozporządzenia stanowią systemy SI wysokiego ryzyka[9]. Do tej kategorii jako pierwsze kwalifikują się systemy SI, które są elementami związanymi z bezpieczeństwem[10] produktów lub same są produktami objętymi wymienionymi w *Aneksie II* unijnymi aktami harmonizującymi i zgodnie z tymi aktami muszą przejść badanie zgodności przez stronę trzecią (celem otrzymania certyfikatu zgodności europejskiej – oznaczenia CE). W tej grupie aktów znalazły się m. in. dyrektywy dotyczące bezpieczeństwa zabawek i wind oraz rozporządzenie ws. kolejek linowych.

„Prawdziwa” lista SI wysokiego ryzyka znalazła się w *Aneksie III*. Wskazane tam systemy zostały uszeregowane wg obszarów zastosowania. Z „automatu” za SI wysokiego ryzyka uznaje się system wykorzystywany w produkcji (lub będący produktem) wymagającym oceny zgodności w myśl przepisów unijnych (tu maszyny rolnicze, lotnictwo, pojazdy, środki ochrony osobistej, maszyny pod ciśnieniem, zabawki, sprzęt medyczny itp.) jak i SI zarządzającą i operującą infrastrukturą krytyczną. Za SI wysokiego ryzyka uznaje się też SI służącą do identyfikacji biometrycznej, będącą komponentem systemów bezpieczeństwa ruchu drogowego i dostaw mediów, służącą kwalifikacji i oceny w edukacji i szkoleniu zawodowym; rekrutacji, relacjach pracodawca-pracownik (ocena pracownicza etc); zarządzającą dostępem do państwowych świadczeń socjalnych, oceną zdolności kredytowej,

zarządzaniem priorytetami ratownictwa (w tym Straży Pożarnej), oceną skłonności do przestępstwa[11] w różnych skalach, podatności na wiktymizację, SI używaną w poligrafach, służącą wykrywaniu tzw. deep fakes i oceną wiarygodności dowodów, wyszukiwanie wzorców przestępstw na dużych wolumenach danych, używaną do zarządzania imigracją, azylem i kontrolą graniczną, a także do wsparcia sądownictwa w wyszukiwaniu i ocenie faktów, prawa i wyrokowaniu (subsumpcji)[12]

Rozporządzenie pozostawia też furtkę Komisji Europejskiej do kwalifikowania SI jako SI wysokiego ryzyka, o ile wykorzystywana jest w jednym z obszarów wskazanych w Aneksie III a równocześnie Komisja uzna, że SI jest wysokiego ryzyka[13].

Nie oznacza to jednak, że każdy system SI używany w obszarach „wysokiego ryzyka” będzie kwalifikował się jako SI wysokiego ryzyka – potrzebne jest również konkretne zastosowanie SI w ramach danego obszaru. W ten sposób systemy SI używane w sektorze pracy są wysokiego ryzyka, jeśli są stosowane przy rekrutacji, przypisywaniu obowiązków, ewaluacji oraz przyznawaniu awansów pracownikom. Podobnie SI stosowana w obszarze wymiaru sprawiedliwości jest wysokiego ryzyka tylko, jeśli służy wspomaganie sędziów w przy wyszukiwaniu i interpretacji faktów oraz stosowania prawa do tych faktów. Szeroko ujęty został jedynie obszar identyfikacji biometrycznej: systemy identyfikacji biometrycznej, zarówno te operujące w czasie rzeczywistym, jak i następczo, będą uznawane za SI wysokiego ryzyka.

Jeżeli natomiast system SI nie kwalifikuje się ani do kategorii wysokiego ryzyka, ani jako w ogóle zakazany, należy przypisać go do trzeciej grupy, którą roboczo można nazwać „SI niskiego ryzyka”. Ryzyko to w ocenie Komisji Europejskiej jest najwyraźniej tak niskie, że tego typu systemy prawie w ogóle nie zostały uregulowane w Rozporządzeniu. Szczątkowych obowiązków można doszukiwać się wśród obowiązków w zakresie przejrzystości, które dotyczą się wszystkich systemów SI, m. in. obowiązku informowania osób wchodzących w interakcje z SI, że mają do czynienia ze Sztuczną Inteligencją.

Jeśli więc w tym momencie można mówić o jakichkolwiek zarzutach dotyczących klasyfikowania systemów SI, raczej nie będą one dotyczyły SI zakazanych czy wysokiego

ryzyka. Te są co prawda wymienione enumeratywnie, ale ich zestaw może się jeszcze wielokrotnie zmienić, ponieważ rozmawiamy jedynie o projekcie. Nawet po wejściu w życie Rozporządzenia, poza generalnym obowiązkiem okresowej kontroli tekstu, w odniesieniu do *Aneksu III*, zawierającego „właściwą” listę SI wysokiego ryzyka, zostały zawarte szeroko opisane kompetencje i kryteria dokonywania zmian. Niepokój budzą natomiast systemy SI „niskiego ryzyka”, wobec których może zdarzyć się, że nie będą objęte żadnymi przepisami. Kiedy obok tego postawimy długą listę obowiązków ciążących na SI wysokiego ryzyka (której najważniejsze elementy zaraz omówimy) obawy budzi ryzyko powstania sytuacji, w których dostawca SI za wszelką cenę będzie chciał zakwalifikować swój produkt do „nieuregulowanej” grupy.

3. Jakie wymagania musi spełnić SI wysokiego ryzyka?

SI wysokiego ryzyka będą musiały być projektowane i tworzone w ścisłej zgodzie z wymaganiami wskazanymi w Rozdziale 2 Rozporządzenia.

Na wymogi względem dostawców SI wysokiego ryzyka składają się: stosowanie systemu zarządzania ryzykiem, zarządzanie danymi (w tym zarządzania uprzedzeniem), opracowywanie i utrzymywanie dokumentacji (technicznej i użytkownika), przechowywanie logów, dbałość o przejrzystość dla użytkownika[14], zapewnienie ludzkiego nadzoru nad działaniem SI, zapewnienie trafności działania SI, odporności (na manipulację, na zmiany danych), zapewnienie cyberbezpieczeństwa SI, stosowanie oceny zgodności, rejestracja SI wysokiego ryzyka, zgłaszanie niezgodności, zapewnienie rozliczalności, stosowanie systemów zarządzania jakością, wyznaczenie reprezentanta w UE, przeprowadzanie oceny skutków dla ochrony danych.

Jeśli system SI wysokiego ryzyka wykorzystuje techniki uczenia się („trenowania”) na danych, to dane te muszą spełniać konkretne wymagania. Wszystkie zestawy danych: „treningowe”, „sprawdzające” oraz „testowe” muszą podlegać odpowiednim praktykom zarządzania.

Ponadto zestawy danych muszą być istotne, reprezentatywne, kompletne i wolne od błędów. W procesie doboru danych nie bez znaczenia pozostaje projektowane zastosowanie SI, ponieważ zestawy danych powinny uwzględniać uwarunkowania geograficzne, behawioralne oraz funkcjonalne, w których będzie działać SI. Wszystkie przedstawione w Rozporządzeniu zasady zarządzania danymi mają jeden główny cel: przeciwdziałanie uprzedzeniom. Eksperci od dłuższego czasu biją na alarm, wskazując jak duże zagrożenie może stanowić „stronnicza” SI[15]. Kluczem do obiektywizmu systemu są odpowiednie dane, co zostało uwzględnione w Rozporządzeniu.

Każdy, kto kiedykolwiek zetknął się z oprogramowaniem, wie, jak duży problem stanowią braki w dokumentacji. W przypadku systemów SI wysokiego ryzyka tego typu problemy są niedopuszczalne (choć równocześnie nieuniknione – podobnie jak przy oprogramowaniu), stąd Rozporządzenie nakłada obowiązek sporządzenia kompletnej dokumentacji jeszcze przed wprowadzeniem SI do obrotu, a także jej bieżącego aktualizowania[16]. Dokumentacja musi jasno wskazywać, że SI wysokiego ryzyka spełnia wymogi Rozporządzenia. Powyższe uzupełnia obowiązek sporządzania i przechowywania logów tak, aby dało się dogłębnie (sic?) prześledzić funkcjonowanie SI.

Kolejny wielokrotnie podnoszony problem Sztucznej Inteligencji dotyczy tak zwanej „czarnej skrzynki”. Samouczące się systemy często tworzą pułapkę, w której dla konkretnych danych SI produkuje (zwykle) prawidłowy wynik, ale nikt nie jest w stanie stwierdzić, w jaki sposób system do tego wyniku doszedł. W związku z tym Rozporządzenie wymaga od systemów SI wysokiego ryzyka zachowania odpowiedniego poziomu przejrzystości. SI powinny być tworzone w taki sposób, aby korzystające z nich podmioty mogły odpowiednio interpretować i wykorzystywać uzyskiwane wyniki. W osiągnięciu tego celu mają pomóc obszerne instrukcje obsługi, a także interfejsy umożliwiające ludzki nadzór nad pracą SI. Okaże się w przyszłości, czy ten dezyderat pozostanie w sferze myślenia życzeniowego czy wywrze wpływ na rynek.

Ciekawe sformułowania pojawiają się przy obowiązkach dotyczących dokładności (*accuracy*), solidności (*robustness*) i cyberbezpieczeństwa. Systemy SI wysokiego ryzyka, które kontynuują naukę po wprowadzeniu do obrotu, muszą być odporne na powstawanie tzw.

feedback loops, czyli sytuacji, w których wyniki produkowane przez SI zostają w dużej mierze oparte o wcześniejsze wyniki dawane przez tą samą SI. W zakresie cyberbezpieczeństwa SI muszą być odporne na *data poisoning*, czyli ataki nakierowane na manipulowanie zestawami danych, na których trenuje SI, oraz na *adversarial examples*, czyli wprowadzanie danych mających na celu spowodowanie pomyłki samouczącego się systemu.

Systemy SI wysokiego ryzyka wskazane w *Aneksie III* (czyli te, których nie pokrywają osobne unijne akty harmonizujące) będą również podlegały obowiązkowej rejestracji w unijnej bazie danych.

Powyższe wymagania, wraz z innymi wymaganiami dotyczącymi SI wysokiego ryzyka, stanowią fundamenty, na których powinien stać każdy system tego typu. W zależności od roli, w jakiej dany podmiot występuje wobec SI, obowiązki te mogą jednak dotyczyć go w różnym stopniu. W ten sposób podmiot stojący za stworzeniem SI (*dostawca*) ma bardzo długą listę obowiązków, dotyczącą ciągłego zarządzania jakością i zgodnością z prawem jego systemu. Natomiast obowiązki podmiotu korzystającego z SI w ramach swojej działalności (*użytkownik*) w dużej mierze ograniczają się do podążania za instrukcją obsługi. Uwaga: jeśli dany podmiot dokona znaczących modyfikacji w systemie SI, w szczególności modyfikacji jego przewidywanego zastosowania, lub jeśli oznaczy SI swoim logo, przestanie być traktowany jedynie jako *użytkownik* i zamiast tego będzie traktowany jako *dostawca*, z wszystkimi obowiązkami, które za tym stoją.

4. Rozporządzenie AI a ochrona danych

Z uwagi na charakter niniejszej publikacji nie sposób pominąć tematu relacji Rozporządzenia z obszarem ochrony danych osobowych. Szczególnie, że pojawia się sporo punktów, w których obie te tematyki się przecinają.

Po pierwsze, powyżej omawialiśmy już kwestię danych używanych do trenowania SI. W celu przeciwdziałania uprzedzeniom zestawy danych muszą być m. in. kompletne, wolne od

błędów czy uwzględniać lokalne uwarunkowania rynku, na którym działać będzie SI. To oznacza, że SI co do zasady powinna być trenowana na prawdziwych danych. A ponieważ rozmawiamy przede wszystkim o zapobieganiu uprzedzeniom, Rozporządzenie wprost zezwala na przetwarzanie w tym celu danych szczególnych kategorii – oczywiście przy uwzględnieniu wszelkich możliwych zabezpieczeń i tylko w granicach uzasadnionych tym celem. Zgodnie z treścią motywów Rozporządzenie co do zasady nie może być jednak traktowane jako dające samodzielną podstawę prawną do przetwarzania danych osobowych. Drugim wyjątkiem w tej kwestii jest korzystanie ze zbiorów danych zebranych w innych celach w ramach *piaskownic regulacyjnych*, mających ułatwić mniejszym graczom wejście na rynek Sztucznej Inteligencji.

Po drugie, projekt Rozporządzenia był omawiany podczas czerwcowych obrad plenarnych Europejskiej Rady Ochrony Danych (EROD), w efekcie czego powstało stanowisko w tej sprawie, sporządzone wspólnie z Europejskim Inspektorem Ochrony Danych[17]. Część propozycji została przez EROD powitana z otwartymi ramionami, w szczególności przyjęcie przez Rozporządzenie podejścia opartego na ryzyku (podobnie jak przy RODO). EROD pozytywnie ocenia również zawarte w Rozporządzeniu zastrzeżenie, że samo zakwalifikowanie systemu SI jako wysokiego ryzyka nie gwarantuje jeszcze, że jego użycie będzie zgodne z prawem – ponieważ istnieją jeszcze inne unijne akty, takie jak RODO, z którymi należy osobno zapewnić zgodność.

W wielu punktach EROD wskazuje na dobre w zamyśle postanowienia Rozporządzenia, które zdaniem Rady wymagają jednak pogłębienia. Chodzi tutaj przede wszystkim o wzmocnienie związku Rozporządzenia z RODO – i to nie tylko na papierze. Zgodnie z propozycją EROD spełnienie wymogów wynikających z prawa unijnego (w tym dotyczących ochrony danych) powinno być warunkiem uzyskania oznaczenia CE. Oprócz powyższego EROD wskazuje wiele obszarów Rozporządzenia – kodeksy postępowania, piaskownice regulacyjne, system certyfikacji – w których zdaniem Rady powinno znaleźć się więcej precyzyjnych odniesień do unijnego prawa ochrony danych (w skrócie: żeby musiały spełniać lub wymagać od zainteresowanych spełniania). Ponieważ Sztuczna Inteligencja wydaje się nierozzerwalnie

połączona z danymi osobowymi (i to na wielu płaszczyznach: od trenowania po wchodzenie w interakcje), a z drugiej strony uregulowania dotyczące ochrony danych muszą być silnie zakorzenione w prawodawstwie, nie sposób nie zgodzić się z postulatami EROD równocześnie uznając je za ogólnikowe.

Konkretne elementy sporne pojawiają się w kontekście technologii mogących mieć zdaniem EROD niekorzystny wpływ na prawa i wolności jednostek. W kwestii systemów *scoringu* społecznego EROD postuluje całkowite ich zakazanie. Jak wskazuje, podmioty prywatne, w szczególności z branży mediów społecznościowych oraz usług chmurowych, przetwarzają ogromne ilości danych osobowych, na podstawie których prowadzą *scoring* społeczny. Nie ma więc powodów, dla których zakaz miałby objąć tylko władze publiczne. Długa lista zarzutów pojawia się również pod adresem systemów identyfikacji biometrycznej. Zdaniem EROD zautomatyzowane rozpoznawanie osób fizycznych powinno w ogóle być w przestrzeniach publicznych zakazane, ponieważ, poza prostą potrzebą pozostania anonimowym, takie działanie SI praktycznie nie może przestrzegać zasad niezbędności i proporcjonalności przetwarzania danych. Zakazana powinna zostać również kategoryzacja biometryczna (czyli grupowanie ludzi ze względu na ich cechy, w szczególności cechy często stanowiące podstawy dyskryminacji), jako że nie odpowiada godności ludzkiej i w krótkiej perspektywie może prowadzić do uprzedmiotowienia osób, których dane dotyczą.

Po trzecie, otwarta pozostaje kwestia organów nadzorujących przestrzeganie Rozporządzenia. Z Rozporządzenia wiemy już, że organem nadzorczym na poziomie unijnym miałby być Europejski Inspektor Ochrony Danych. Biorąc pod uwagę to, postulaty EROD w tym zakresie oraz mnogość obszarów, w których tematyka SI przecina się z tematyką ochrony danych osobowych, wydaje się, że na poziomie krajowym rolę organów nadzorczych do spraw SI powinny objąć istniejące już organy nadzorcze do spraw ochrony danych osobowych. Takie podejście jest słuszne tym bardziej dlatego, że krajowe organy nadzorcze ds. ochrony danych już teraz egzekwują przestrzeganie odpowiednich przepisów przez działające już systemy SI, które korzystają z danych osobowych. Co za tym idzie organy te zdążyły już osiąść pewną wiedzę na temat działania tych technologii.

5. Perspektywy na przyszłość

W niniejszej publikacji staraliśmy się pokryć najistotniejsze elementy Rozporządzenia – opisać kierunki, w których prawdopodobnie pójdzie regulacja SI. Projekt jest obszerny, a w najbliższym czasie niewątpliwie będzie podlegać szerokim unijnym uzgodnieniom, a następnie zmianom wynikającym z tych uzgodnień. Pewne wątpliwości można jednak przedstawić już teraz.

Samo zainteresowanie się tematyką SI przez unijnego ustawodawcę należy ocenić jednoznacznie pozytywnie. Regulacja tego obszaru posłuży wszystkim uczestnikom obrotu gospodarczego. Najbardziej oczywiście osobom fizycznym, ponieważ reprezentowana przez Rozporządzenie idea *human centered AI* stawia ich prawa i wolności na pierwszym miejscu. Jednak w naszej opinii w ostatecznym rozrachunku regulacje będzie korzystna również dla pozostałych uczestników rynku – twórców SI czy podmiotów korzystających z SI w ramach swojej działalności – ponieważ upowszechni dostęp do Sztucznej Inteligencji oraz ustali uczciwe ramy konkurencji. Specjalną uwagę poświęcono nawet małym przedsiębiorcom, którzy będą mieć priorytetowy dostęp do *piaskownic regulacyjnych* tak, aby ich wielkość nie stawiała ich w gorszej pozycji względem większych podmiotów.

Dziwi natomiast fragmentaryczne podejście do tematu. Zrozumiałym jest, że uczestnicy rynku zawsze będą starali się tworzyć swoje produkty tak, aby dotyczyło ich jak najmniej regulacji. Jednak kiedy część systemów SI otrzymuje bardzo szeroki zestaw obowiązków, a pozostałej części prawie żadne obowiązki nie dotyczą, obawy budzi ryzyko powstania sytuacji, w których SI wysokiego ryzyka zostaje przez dostawcę przedstawiona jako SI niskiego ryzyka i tym samym wymyka się regulacjom. Tego typu działanie nie musi nawet być konsekwencją złej woli dostawcy, ponieważ systemy SI wysokiego ryzyka są mniej lub bardziej enumeratywnie wskazane w zamkniętym katalogu. Tym samym albo przez błąd ustawodawcy, albo (co bardziej prawdopodobne) przez standardowo dynamiczny rozwój technologii mogą powstawać SI stanowiące ryzyko dla praw i wolności obywateli, na których równocześnie nie będą ciążyć żadne obowiązki.

Natomiast podążając za wątpliwościami przedstawianymi przez EROD, czeka nas poważna dyskusja na temat granic ochrony prywatności. Niektóre z zastosowań SI zakładają przetwarzanie danych osobowych, w tym zestawianie różnych zbiorów danych, w hurtowych ilościach, praktycznie bez wiedzy osób, których dane dotyczą. Takie systemy – *scoringu* społecznego czy identyfikacji lub kategoryzacji biometrycznej – prawdopodobnie mogą nieść za sobą jakieś korzyści dla osób, których dane przetwarzają (nawet, jeśli w tym momencie autorzy żadnych takich korzyści nie dostrzegają). Z drugiej strony ryzyko, jakie niesie za sobą pozostawienie tak ogromnych zbiorów danych do przetwarzania przez maszyny, które z założenia będą pracować na tych zbiorach z niewielkim udziałem czynnika ludzkiego, może okazać się zbyt wysokie, by było dopuszczalne. A ponieważ, jak już opisywaliśmy powyżej, Rozporządzenie stawia cienkie granice pomiędzy SI niedopuszczalną do stosowania, SI regulowaną oraz SI nieregulowaną, być może rozsądnym wyjściem jest przezorne zakwalifikowanie tak „groźnych” SI do grupy zakazanej.

Zastrzeżenia budzi również projektowane podejście do implementacji Rozporządzenia. Projekt zakłada upływ 24 miesięcy od wejścia w życie do rozpoczęcia stosowania Rozporządzenia. Taki termin jest w pewien sposób zrozumiały, w końcu trzeba nie tylko dać zainteresowanym podmiotom czas na uwzględnienie nowych przepisów, ale przede wszystkim zbudować całą europejską sieć instytucji. Co ważniejsze, projekt zakłada, że do SI wysokiego ryzyka wprowadzonych do obrotu przed datą rozpoczęcia stosowania Rozporządzenia nowe przepisy będą się stosowały tylko, jeśli od tamtego momentu SI ulegnie istotnym zmianom w konstrukcji lub zakładanym zastosowaniu. Takie podejście ustawodawcy oceniamy raczej negatywnie. Nie widzimy powodu, dla którego nie należy zarządzić ogólnego przeglądu systemów SI aktualnie stosowanych na terenie UE. Tym bardziej nie widzimy powodu, dla którego SI wprowadzone do obrotu pomiędzy datą wejścia w życie rozporządzenia (a więc po momencie, w którym ostateczna treść regulacji zostanie ustalona), ale przed rozpoczęciem stosowania Rozporządzenia, miałyby nie podlegać nowym przepisom.

O ile więc pewne – miejscami bardzo istotne – kwestie pozostają otwarte, pierwsze podejście unijnego ustawodawcy do uregulowania Sztucznej Inteligencji należy raczej oceniać pozytywnie. Rozporządzenie przynajmniej stara się podążać za ideą *human centered AI* oraz być tworzone z myślą o dynamicznym rozwoju tego typu technologii. Tak długo, jak ustawodawca będzie trzymał się tych założeń, pozostajemy ostrożnie optymistyczni.

[1] Na potrzeby niniejszej publikacji nazywane po prostu *Rozporządzeniem*. Pełny tytuł polski: 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. Pełny tytuł angielski: Proposal for a Regulation Of The European Parliament And Of The Council Laying Down Harmonised Rules On Artificial Intelligence (Artificial Intelligence Act) And Amending Certain Union Legislative Acts COM/2021/206 final.

[2] Dla większości z Państwa będzie to zapewne Terminator (czy bardziej precyzyjnie system Skynet).

[3] A jeśli ktoś jest dalej zainteresowany temat obywatelstwa Sophii, polecamy artykuł z portalu slate.com zatytułowany *What Exactly Does It Mean to Give a Robot Citizenship?* (Co dokładnie oznacza nadanie robotowi obywatelstwa?):

<https://slate.com/technology/2017/11/what-rights-does-a-robot-get-with-citizenship.html>

[4] *The science of making machines do things that would require intelligence if done by men* (<https://www.britannica.com/biography/Marvin-Lee-Minsky> dostęp 23.07.21)

[5] <https://www.theguardian.com/technology/2021/jun/06/microsofts-kate-crawford-ai-is-neither-artificial-nor-intelligent>

[6] *prohibited Artificial Intelligence practices*

[7] *subliminal techniques*

[8] Tłumaczenie jest nieidealne, ponieważ słowo *xinyong* oznacza zarówno kredyt zaufania, jak i kredyt finansowy. W związku z tym oraz ponieważ zakładają one istnienie mniej lub bardziej rzeczywistych „punktów”, będziemy określać je jako systemy *scoringu* społecznego. (<https://instytutprawobywatelskich.pl/chinski-system-oceny-spoleczenstwa-orwell-i-huxley-w-praktyce/> dostęp 23.07.21)

[9] *high-risk AI systems*

[10] *safety component*

[11] Przypominamy sobie słynny „Raport Mniejszości”. Jednak już teraz funkcjonują na świecie systemy oceniające np. skłonność do recydywy – tu amerykański system COMPAS (Correctional Offender Management Profiling for Alternative Sanctions).

[12] Nie wiemy, na jakim etapie jest estoński projekt rozsądzania małych spraw przez SI ale w Kanadzie podobno robot już przeprowadził skutecznie mediację.

[13] W Rozporządzeniu jest to napisane w sposób nieco zawoalowany, ale taki jest sens.

[14] W branży SI funkcjonuje taki żart, że „algorytm = programista nie chce powiedzieć, co zrobił; heurystyka = programista nie może wyjaśnić, co zrobił; uczenie maszynowe = programista nie wie, co zrobił”. Stąd przejrzystość dla użytkownika była i pozostanie ciekawym wyzwaniem dla twórców sztucznej inteligencji.

[15] <https://www.forbes.com/sites/forbestechcouncil/2021/02/04/the-role-of-bias-in-artificial-intelligence/> (dostęp 23.07.21)

[16] Spodziewamy się (może myśląc nieco cynicznie, a może po prostu w oparciu o doświadczenie życiowe), że powstaną gotowe wzorce dokumentacji SI wysokiego ryzyka, niekoniecznie odzwierciedlające rzeczywiste algorytmy, ale odpowiadające na wyzwania

regulacyjne wprowadzane przez Rozporządzenie. Można rozsądnie oczekiwać, że prawnicy będą brać udział w tworzeniu takiej dokumentacji.

[17] https://edpb.europa.eu/our-work-tools/our-documents/edpbedps-joint-opinion/edpb-edps-joint-opinion-52021-proposal_en (dostęp 26.07.21)



Obiektywnie czy subiektywnie? Jak patrzeć na dane?

Autor: Michał Lewandowski

Cloudtechnologies.pl

W kontekście ustalania czy dany zbiór danych będzie posiadał rangę danych osobowych, zgodnie z rozumieniem przepisów o ochronie danych osobowych kluczową rolę odgrywa ustalenie czy na podstawie zebranych danych możliwa będzie następcza identyfikacja osoby, której dane dotyczą. Kwestie związane z "identyfikowalnością" podmiotu danych osobowych nie uregulowano natomiast w tekście głównym RODO, a jedynie wskazano w motywie 26 RODO, zgodnie z którym aby stwierdzić czy dana osoba jest możliwa do zidentyfikowania należy uwzględnić wszelkie "rozsądnie prawdopodobne sposoby (w tym wyodrębnienie wpisów dotyczących tej samej osoby), w stosunku do których istnieje uzasadnione prawdopodobieństwo, iż zostaną wykorzystane przez administratora lub inną osobę w celu bezpośredniego lub pośredniego zidentyfikowania osoby fizycznej". Ustalając to należy uwzględnić takie czynniki jak koszt i czas niezbędny do zidentyfikowania danej osoby, a także dostępną w danym czasie technologię oraz sam postęp techniczny.

W tym kontekście kolejną istotną rolę odgrywa ustalenie czy na przesłankę identyfikowalności należy patrzeć subiektywnie czy też obiektywnie, tj. czy identyfikowalność należy ustalać z perspektywy wyłącznie podmiotu agregującego dane, czy też w oparciu o pełne spektrum możliwości ich następczego "zderzenia" ze zbiorami innych administratorów danych i następczej identyfikacji podmiotu danych na tej podstawie.

Kwestia ta została poruszona w głośnym wyroku Trybunału Sprawiedliwości UE w sprawie Patrick Breyer (C-582/14) i załączonej do tego wyroku opinii Rzecznika Generalnego M. Camposa Sáncheza-Bordony. W przedstawionym w załączonej do wyroku opinii rozumowaniu opowiedziano się wyraźnie za ujęciem obiektywnym i wskazano, że o możliwości identyfikacji przesądza możliwość podmiotu, który tej identyfikacji ma dokonać, do ustalenia danych osoby fizycznej poprzez "połączenie [danych posiadanych] z danymi dostarczonymi przez osobę trzecią". Zgodnie z ujęciem subiektywnym o kwalifikowalności danego zbioru danych jako danych osobowych decydujące byłoby ustalenie czy podmiot dysponujący danymi może ustalić tożsamość danej osoby (choćby pośrednio) bazując wyłącznie na danych, które sam posiada.

Wskazać należy, że każda z przedmiotowych koncepcji ma swoje wady i zalety. Ujęcie obiektywne zapewnia większą ochronę danych bowiem zapewnia do nich stosowanie wszelkich konsekwencji przewidzianych RODO. Niemniej trudno nie dostrzec, że wskutek braku wyraźnej regulacji prawnej administratorzy danych mogą często być postawieni w sytuacji swoistej niepewności co do tego czy dane zbiory danych należy traktować już jako dane osobowe, czy też jeszcze jako "zwykłe" dane. Ujęcie subiektywne zapewnia natomiast pewność co do charakteru danego zbioru wobec podmiotu, który go przetwarza, jednak, przynajmniej potencjalnie, może wpływać na zmniejszenie zakresu bezpieczeństwa przetwarzania takich danych.

W konkluzji wskazanego wyroku Trybunał opowiada za stanowiskiem obiektywnym uznając dynamiczny adres IP za dane osobowe w odniesieniu do dostawcy usług internetowych, "ze względu na istnienie osoby trzeciej (dostawcy dostępu do sieci), do której można się zwrócić w celu uzyskania innych danych dodatkowych umożliwiających w połączeniu tym

z adresem identyfikację użytkownika". Trybunał uznał przy tym, że do możliwości ustalenia danych osoby fizycznej zachodzi w sytuacji, gdy ich ustalenie na podstawie informacji ze zbiorów dwóch podmiotów jest prawnie dozwolone.

W miejscu warto odwołać się do ustaleń orzecznictwa polskiego wydanego w kontekście charakteru prawnego danych określających numery rejestracyjne pojazdu, które zdaje się zmierzać zupełnie inną linią. Zgodnie z wyrokiem NSA z dnia 11 kwietnia 2019 r., sygn. akt I OSK 1240/17, "numerowi rejestracyjnemu pojazdowi nie można przypisać statusu danych osobowych. W przypadku tradycyjnego numeru rejestracyjnego pojazdu, złożonego z liter i cyfr, określenie tożsamości konkretnej osoby, czy to właściciela, czy też posiadacza pojazdu, wymagałoby bowiem nadmiernych kosztów, czasu i działań". Bardzo podobny pogląd zaprezentowano w wyroku WSA w Gliwicach z dnia 31 października 2018 r., sygn. akt II SA/GL 593/17. Abstrahując od podnoszonej przesłanki koniecznych do poniesienia nakładów, która wraz z wejściem w życie RODO straciła na znaczeniu, ustalenia przedmiotowych wyroków zdają się słuszne z tym zastrzeżeniem, że należy je uzupełnić o wskazanie, że podmioty agregujące przedmiotowe dane powinni móc wykazać, że nie są również w stanie w sposób legalny (poza szczególnymi uprawnieniami prawnymi, np. w związku z dochodzeniem roszczeń) wejść w posiadanie dodatkowych danych o użytkownikach pojazdów od innych administratorów danych. Zastrzec należy, że powyższe orzeczenia nie stanowią jednolitej linii orzeczniczej. Ponadto orzeczenia te nie zapadły w sporach z udziałem PUODO czy też GIODO co osłabia niejako ich rangę.

Niezależnie od powyższego takie ujęcie obiektywnego podejścia do ustalania danych osobowych daje z jednej strony zabezpieczenie prawnego zabezpieczenia tych danych, a z drugiej zapewnia podmiotom administrującym możliwość ścisłej weryfikacji czy występują w roli administratora danych osobowych. Ma ono jednak pewne negatywne oddziaływanie. Rozciąga ono bowiem wszelkie konsekwencje przewidziane wymogami

regulacji ochrony danych osobowych również na te podmioty, które choć teoretycznie mają możliwość wzbogacenia administrowanych baz o dane umożliwiające identyfikację osób nie tylko tego nie podejmują, ale również w ramach prowadzonej działalności podejmują wszelkie wysiłki aby dane te nie mogły posłużyć do następcej identyfikacji. Ścisła orientacja na rzecz podejścia obiektywnego kosztem koncepcji subiektywnej ustalania danych nie wydaje się zatem słuszna. W ocenie autora w przestrzeni podmiotów agregujących dane występuje coraz bardziej widoczna potrzeba ponownego rozpatrzenia przedmiotowej kwestii i ujęcia w formie legislacyjnej.

Zaznaczenia wymaga, że przedmiotowy problem może zostać następczo "ominięty" za sprawą rozwoju technologii kryptograficznych. Nietrudno wyobrazić sobie bowiem rozwiązanie oparte na technologii blockchain, które będzie przypisywać szyfrowany kryptograficznie identyfikator, do którego następczo będą przypisywane inne informacje (niepozwalające na identyfikację) dotyczące danej osoby. W takim rozwiązaniu sam operator technologii będzie mimo wszystko uznawany za administratora danych osobowych, ale jego użytkownicy wskutek technicznej i obiektywnej niemożności rozszyfrowania informacji pozwalających na identyfikację podmiotu danych już nie. Czas pokaże zatem czy rozwój technologii zdezaktualizuje zaprezentowany powyżej problem.



URZĄD OCHRONY DANYCH OSOBOWYCH