

Stany Zjednoczone regulują sztuczną inteligencję

11 listopada 2023

Pierwsza regulacja AI stała się faktem. Nie za sprawą głośnego w Unii Europejskiej aktu o sztucznej inteligencji, ale dekretu, który podpisał właśnie prezydent Stanów Zjednoczonych. To jednak przede wszystkim polityczna deklaracja. Pierwsza kompleksowa regulacja ma powstać jeszcze w tym roku w Unii Europejskiej. Jak dobrze ochroni ludzi i ich prawa?

Rozporządzenie w sprawie bezpiecznej i godnej zaufania sztucznej inteligencji (Order on Safe, Secure, and Trustworthy Artificial Intelligence) ma „wyznaczać nowe standardy dla bezpieczeństwa AI, chronić prywatność Amerykanów i Amerykanek, promować równość i prawa obywatelskie, chronić konsumentów i konsumentki oraz pracowników i pracowniczki, promować innowacje i konkurencyjność”, a także „umacniać amerykańskie przywództwo na świecie”.

Kiedy 30 października 2023 r. Joe Biden podpisywał swój dekret, w Unii Europejskiej nadal toczyły się (i wciąż się toczą) negocjacje trójstronne wokół rozporządzenia mającego regulować sztuczną inteligencję na Starym Kontynencie. Czy to znaczy, że AI Act nie będzie „pierwszą kompleksową regulacją AI na świecie”? Wciąż ma szansę w tym wyścigu, bo to dokumenty o zupełnie innej randze i ambicjach.

Rozporządzenie wykonawcze Bidena wspomina wszystkie głośne problemy ze sztuczną inteligencją:

- zapowiada walkę z deepfake’ami i dezinformacją, m.in. poprzez wprowadzenie zabezpieczeń oficjalnych rządowych materiałów, tak aby można je było odróżnić od tych wyprodukowanych przez AI;

- ma chronić Amerykanów i Amerykanki przed dyskryminacją, m.in. na rynku nieruchomości, w dostępie do edukacji czy pracy, a także w zetknięciu z wymiarem sprawiedliwości (pamiętacie jeszcze niesławny COMPAS, który miał wspierać działanie systemu sprawiedliwości, ale dyskryminował osoby inne niż biali mężczyźni w średnim wieku, zawyżając ryzyko, że ponownie złamią prawo?);
- zapowiada też wprowadzenie zasad dotyczących tworzenia, kupowania i używania systemów AI przez wojsko i służby.

Ważniejsze jest jednak to, czego w dekreście nie ma, albo raczej, na czym polega jego ograniczona moc.

Rozporządzenie wykonawcze prezydenta USA to przede wszystkim polityczna deklaracja. Pokazuje kierunek dalszych działań, ale nie ma rangi ustawy, więc nie rozwiąże fundamentalnych problemów. Takim problemem jest na przykład brak kompleksowego prawa chroniącego prywatność wszystkich Amerykanów i Amerykanek. W dekreście prezydent jedynie apeluje do Kongresu o jego uchwalenie. W innych kwestiach Biden nakazuje podległym mu instytucjom podjąć konkretne działania. Nie wiadomo, jak te zmiany będą wyglądać w praktyce. Moc dekretu może też nie wystarczyć w starciach sądowych, jeśli firma produkująca AI pozwie administrację za działania, które uzna za naruszające amerykańskie prawo. Co więcej: dekret może też łatwo odwołać sam prezydent. Jeśli nie obecny, to jego następca lub następczyni.

Unijne rozporządzenie nie ma aż tak wielkich ambicji jak dokument podpisany przez Bidena. Różni się zakresem: reguluje „tylko” rynek, pozostawiając z boku wykorzystanie AI przez wojsko krajów członkowskich. Ma natomiast rangę prawa obowiązującego bezpośrednio we wszystkich państwach UE. Wprowadza też bardzo szczegółowe obostrzenia: w tym listę zakazanych systemów i liczne wymogi wobec systemów z kategorii tzw. wysokiego ryzyka. Wszystkie te rozwiązania mają chronić ludzi i ich prawa przed zagrażającym im wykorzystaniem

systemów AI. To, jak mocna będzie ta ochrona, zależy od ostatecznego kształtu regulacji, które zostaną przyjęte w ramach negocjacji trójstronnych.

To jednak pole przeciągania liny między różnymi grupami interesów. Organizacje społeczne, w tym Fundacja Panoptikon, i Parlament Europejski chciałyby większej ochrony. Natomiast niektóre państwa zrzeszone w Radzie Unii Europejskiej i big techy dążą do większej swobody w wykorzystaniu narzędzi AI, zwłaszcza pod pretekstem bezpieczeństwa lub rozwoju innowacji. Szczególną areną zmagania są regulacje dotyczące systemów wysokiego ryzyka – przede wszystkim to, jakie systemy załapią się na tę etykietę.

Jeśli wygra interes firm, spod rygoru przewidzianego dla systemów wysokiego ryzyka wymkną się np. takie narzędzia, których celem jest „ulepszenie wyniku działania wykonanego wcześniej przez człowieka” (zgodzicie się, że to dość pojemna kategoria...). Chyba że prowadzą do profilowania osób fizycznych: strony porozumiały się, że takie systemy będą – bez wyjątku – traktowane jako systemy wysokiego ryzyka, wraz ze wszystkimi konsekwencjami przewidzianymi dla tej kategorii, na czele z wymogiem przeprowadzenia oceny skutków dla praw podstawowych i wpisania do centralnej bazy.

Rządy zaś dążą do wprowadzenia w rozporządzeniu wyjątków wobec systemów używanych w polityce migracyjnej czy do biometrycznej inwigilacji w czasie rzeczywistym. Możliwe luzowanie wymogów wobec nich byłoby złą wiadomością dla obywateli i obywaterek, którzy nie mogliby uzyskać podstawowych informacji o narzędziach sztucznej inteligencji wykorzystywanych przez policję czy straż graniczną.

Krytyczny głos na temat kierunku, jaki przybierają na ostatniej prostej negocjacje aktu o AI, zabrał ostatnio Europejski Inspektor Danych Osobowych. Prof. Wojciech Wiewiórowski jasno opowiedział się za ochroną ludzi. W swojej opinii podkreślił, że zakres zakazanych systemów AI jest zbyt

wąski, w systemie klasyfikowania systemów jako AI wysokiego ryzyka są luki, a ostateczna wersja aktu powinna przyznawać osobom poddanym działaniu systemów AI konkretne prawa, np. odwołania się od decyzji podjętej z wykorzystaniem takiego systemu.

W czerwcu chwalono Parlament Europejski za wciągnięcie na listę systemów wysokiego ryzyka algorytmów rekomendacyjnych wielkich platform internetowych (VL0P-ów). Negocjacje zmierzają do wykreślenia ich z tej listy. Nawet gdyby się tak stało, wciąż mogłyby podlegać wymogom dla takich systemów ze względu na to, że ich działanie stanowi „znaczące” ryzyko szkód dla zdrowia, życia lub praw podstawowych osób poddanych ich działaniu.

Big techy jednak robią, co mogą, żeby uniknąć kłopotliwej dla nich etykiety. Meta na przykład dowodzi, że ich rekomendacje nie mogą być aż tak szkodliwe, jak dajmy na to (błędne) działanie systemu na użytek wymiaru sprawiedliwości. Taki argument padł przynajmniej w przeprowadzonych niedawno przez polski rząd konsultacjach (zobacz odpowiedzi, jakie przesłały firmy). Platforma zdaje się ignorować głośne przypadki samobójstw wśród młodych osób, które wpadły w „króliczą norę” rekomendacji na „Facebooku” czy „Instagramie”. Chociaż trudno zero-jedynkowo wykazać, że winne są właśnie rekomendacje, to udowodniono już, że serwowanie przez Metę czy „TikToka” treści związanych z samookaleczeniem, gloryfikujących samobójstwo czy napędzających zaburzenia odżywiania może pogłębiać problemy psychiczne.

Cieszono się też ze wpisania przez Parlament systemów służących do identyfikacji biometrycznej w czasie rzeczywistym na listę systemów zakazanych. Miano nadzieję, że – dzięki temu, że rozporządzenie UE ma wyższą rangę niż ustawa krajowa – zneutralizuje to przyjęte przez Francję przepisy dopuszczające takie rozwiązanie w związku z organizacją igrzysk olimpijskich i paraolimpijskich w Paryżu w 2024 r. To się nie wydarzy, i to nie tylko dlatego, że igrzyska już w

przyszłym roku, a AI Act zacznie obowiązywać dopiero za wiele miesięcy, ale też dlatego, że nie wiadomo, czy ten zakaz się utrzyma w ostatecznej wersji rozporządzenia.

To zresztą m.in. francuski rząd dąży do tego, żeby całkowity zakaz systemów biometrycznych został zastąpiony konkretnymi wymogami, ograniczającymi możliwość ich użycia do „wyjątkowych sytuacji” i „najpoważniejszych przestępstw”. Ciekawe, czy pojawią się jakieś nowe preteksty równie pojemne, co wyświechtane już „zwalczanie terroryzmu” czy „bezpieczeństwo narodowe”.

Czy unijne instytucje w kończących się właśnie negocjacjach trójstronnych oprą się pokusie rozszczelnienia (dość dobrej) ochrony dla ludzi i praw, którą udało się stworzyć w wersji aktu o AI przyjętej przez Parlament Europejski?

Autorstwo: Anna Obem

Współpraca: Filip Konopczyński, Maria Wróblewska

Źródło: [Panoptykon.org](https://panoptykon.org)