



KOMISJA
EUROPEJSKA

Bruksela, dnia 21.4.2021 r.
SWD(2021) 85 final

DOKUMENT ROBOCZY SŁUŻB KOMISJI
STRESZCZENIE SPRAWOZDANIA Z OCENY SKUTKÓW

Towarzyszący dokumentowi:

Wniosek dotyczący rozporządzenia Parlamentu Europejskiego i Rady
USTANAWIAJĄCEGO ZHARMONIZOWANE PRZEPISY DOTYCZĄCE
SZTUCZNEJ INTELIGENCJI (AKT W SPRAWIE SZTUCZNEJ INTELIGENCJI) I
ZMIENIAJĄCEGO NIEKTÓRE AKTY USTAWODAWCZE UNII

{COM(2021) 206 final} - {SEC(2021) 167 final} - {SWD(2021) 84 final}

Streszczenie oceny skutków
Ocena skutków ram regulacyjnych w zakresie sztucznej inteligencji
A. Zasadność działań
Na czym polega problem i dlaczego jest to problem na szczeblu UE?
Sztuczna inteligencja (AI) jest powstającą technologią o szerokim zakresie zastosowań: jest to zbiór bardzo potężnych technik programowania komputerowego. Upowszechnienie systemów sztucznej inteligencji niesie ze sobą duży potencjał: może zapewnić korzyści społeczne, wzrost gospodarczy oraz pobudzić innowacyjność i zwiększyć globalną konkurencyjność UE. W niektórych przypadkach wykorzystywanie systemów sztucznej inteligencji może jednak stwarzać problemy. Szczególne cechy niektórych systemów sztucznej inteligencji mogą spowodować nowe zagrożenia związane z 1) bezpieczeństwem i ochroną oraz 2) prawami podstawowymi, a także zwiększyć prawdopodobieństwo wystąpienia lub intensywność istniejących zagrożeń. Systemy sztucznej inteligencji 3) utrudniają również organom ścigania weryfikację zgodności z obowiązującymi przepisami i ich egzekwowanie. Ten zbiór problemów prowadzi z kolei do 4) niepewności prawa, z którą borykają się przedsiębiorstwa, 5) potencjalnie wolniejszego wdrażania technologii AI ze względu na brak zaufania ze strony przedsiębiorstw i obywateli, a także 6) działań regulacyjnych ze strony organów krajowych w celu złagodzenia ewentualnych efektów zewnętrznych, co grozi fragmentacją rynku wewnętrznego.
Co należy osiągnąć?
Ramy regulacyjne służą rozwiązaniu tych problemów w celu zapewnienia prawidłowego funkcjonowania jednolitego rynku przez stworzenie warunków sprzyjających rozwijaniu i wykorzystywaniu w Unii wiarygodnej sztucznej inteligencji. Celami szczegółowymi są: 1) zapewnienie, aby systemy sztucznej inteligencji wprowadzane do obrotu i znajdujące się w użyciu były bezpieczne i zgodne z obowiązującym prawem w obszarze praw podstawowych oraz z unijnymi wartościami; 2) zapewnienie pewności prawa, aby ułatwić inwestycje i innowacje w dziedzinie sztucznej inteligencji; 3) poprawa zarządzania i skuteczne egzekwowanie obowiązujących przepisów dotyczących praw podstawowych i wymogów bezpieczeństwa mających zastosowanie do systemów sztucznej inteligencji oraz 4) ułatwienie rozwoju jednolitego rynku zgodnych z prawem, bezpiecznych i wiarygodnych systemów sztucznej inteligencji oraz zapobieganie fragmentacji rynku.
Na czym polega wartość dodana podjęcia działań na poziomie UE (zasada pomocniczości)?
Transgraniczny charakter wielkoskalowych danych i dużych zbiorów danych, na których często opierają się zastosowania AI, oznacza, że cele inicjatywy nie mogą zostać skutecznie osiągnięte przez same państwa członkowskie. Celem europejskich ram regulacyjnych w zakresie wiarygodnej sztucznej inteligencji jest ustanowienie zharmonizowanych przepisów dotyczących rozwoju, wprowadzania do obrotu i wykorzystywania w Unii produktów i usług z wbudowaną technologią sztucznej inteligencji lub samodzielnych zastosowań AI. Mają one zapewnić równe szanse i ochronę wszystkich obywateli Unii, wzmacniając jednocześnie konkurencyjność i bazę przemysłową Europy w dziedzinie sztucznej inteligencji. Działania UE w zakresie sztucznej inteligencji wzmocnią rynek wewnętrzny i mają znaczny potencjał zapewnienia przemysłowi europejskiemu przewagi konkurencyjnej na poziomie globalnym dzięki korzyściom skali, których nie mogą osiągnąć samodzielnie poszczególne państwa członkowskie.
B. Rozwiązania
Jakie są różne warianty działań służących osiągnięciu celów? Czy wskazano preferowany wariant? Jeżeli nie, dlaczego?
Rozważono następujące warianty: wariant 1 : instrument legislacyjny UE ustanawiający dobrowolny system etykietowania; wariant 2 : podejście sektorowe ad hoc; wariant 3 : horyzontalny instrument legislacyjny UE ustanawiający obowiązkowe wymogi dotyczące zastosowań AI wysokiego ryzyka ; wariant 3+ : taki sam jak wariant 3, ale z dobrowolnymi kodeksami postępowania dla zastosowań AI nieobarczonych wysokim ryzykiem; oraz wariant 4 : horyzontalny instrument legislacyjny UE ustanawiający obowiązkowe wymogi dotyczące wszystkich zastosowań sztucznej inteligencji.

Preferowanym wariantem jest wariant 3+, ponieważ oferuje on proporcjonalne zabezpieczenia przed zagrożeniami stwarzanymi przez sztuczną inteligencję, jednocześnie ograniczając do minimum koszty administracyjne i koszty przestrzegania przepisów. Szczegółowa kwestia odpowiedzialności za zastosowania AI zostanie uregulowana w przyszłości odrębnymi przepisami, a zatem nie uwzględniono jej w poszczególnych wariantach.
Jakie są opinie poszczególnych zainteresowanych stron? Jak kształtuje się poparcie dla poszczególnych wariantów?
Przedsiębiorstwa, organy publiczne, naukowcy i organizacje pozarządowe zgadzają się, że istnieje luka prawna lub że konieczne jest wprowadzenie nowego prawodawstwa, chociaż wśród przedsiębiorstw respondenci zgadzający się takim wnioskiem nie stanowią tak znacznej większości jak w przypadku innych grup respondentów. Przemysł i organy publiczne opowiadają się za ograniczeniem obowiązkowych wymogów do zastosowań AI wysokiego ryzyka. Na ograniczenie obowiązkowych wymogów do zastosowań wysokiego ryzyka częściej nie zgadzają się natomiast obywatele i społeczeństwo obywatelskie.
C. Skutki wdrożenia preferowanego wariantu
Jakie korzyści przyniesie wdrożenie preferowanego wariantu lub – jeśli go nie wskazano – głównych wariantów?
W przypadku obywateli preferowany wariant zmniejszy ryzyko zagrożeń dla ich bezpieczeństwa i praw podstawowych. W przypadku dostawców stworzy on pewność prawa i zagwarantuje, że nie pojawią się żadne przeszkody w transgranicznym świadczeniu usług i oferowaniu produktów związanych ze sztuczną inteligencją. W przypadku przedsiębiorstw korzystających ze sztucznej inteligencji będzie promować zaufanie wśród klientów. W przypadku krajowych organów administracji publicznej będzie promować zaufanie publiczne do wykorzystywania sztucznej inteligencji i wzmocni mechanizmy egzekwowania prawa (poprzez wprowadzenie europejskiego mechanizmu koordynacji, zapewnienie odpowiedniego potencjału i ułatwienie kontroli systemów sztucznej inteligencji dzięki wprowadzeniu nowych wymogów w zakresie dokumentacji, identyfikowalności i przejrzystości).
Jakie są koszty wdrożenia preferowanego wariantu lub – jeśli go nie wskazano – głównych wariantów?
Przedsiębiorstwa lub organy publiczne, które opracowują lub wykorzystują systemy AI stwarzające wysokie ryzyko dla bezpieczeństwa lub praw podstawowych obywateli, będą musiały spełnić szczególne wymogi i obowiązki horyzontalne, które zostaną wprowadzone w życie za pomocą technicznych norm zharmonizowanych. Szacuje się, że do 2025 r. całkowity łączny koszt przestrzegania przepisów wyniesie 100–500 mln EUR, co odpowiada maksimum 4–5 % inwestycji w sztuczną inteligencję wysokiego ryzyka (szacuje się, że jest to 5–15 % wszystkich zastosowań AI). Koszty weryfikacji mogłyby wynieść kolejne 2–5 % w przypadku inwestycji w sztuczną inteligencję wysokiego ryzyka. Przedsiębiorstwa lub organy publiczne, które opracowują lub wykorzystują systemy sztucznej inteligencji niesklasyfikowane jako systemy wysokiego ryzyka, nie musiałyby ponosić żadnych kosztów. Mogłyby jednak zdecydować się na stosowanie się do dobrowolnych kodeksów postępowania, co wiązałoby się z przestrzeganiem odpowiednich wymogów, aby dać pewność, że ich systemy sztucznej inteligencji są wiarygodne. W takich przypadkach koszty mogłyby być co najwyżej równe kosztom ponoszonym w przypadku zastosowań wysokiego ryzyka, ale najprawdopodobniej byłyby niższe.
Jakie są skutki dla MŚP i konkurencyjności?
MŚP odniosą większe korzyści z wyższego ogólnego poziomu zaufania do sztucznej inteligencji niż duże przedsiębiorstwa, które mogą polegać również na wizerunku swojej marki. MŚP rozwijające zastosowania sklasyfikowane jako zastosowania wysokiego ryzyka musiałyby ponosić podobne koszty jak duże przedsiębiorstwa. Ze względu na wysoką skalowalność technologii cyfrowych małe i średnie przedsiębiorstwa mogą mieć ogromny zasięg pomimo swoich niewielkich rozmiarów, potencjalnie wpływając na miliony osób. Dlatego też, jeśli chodzi o zastosowania wysokiego ryzyka, wyłączenie MŚP dostarczających systemy sztucznej inteligencji z zakresu stosowania ram regulacyjnych mogłoby

poważnie zaszkodzić celowi, jakim jest zwiększenie zaufania. W omawianych ramach zostaną jednak przewidziane szczególne środki, w tym piaskownice regulacyjne lub pomoc za pośrednictwem ośrodków innowacji cyfrowych, których celem jest wspieranie MŚP w przestrzeganiu nowych przepisów z uwzględnieniem ich szczególnych potrzeb.
Czy przewiduje się znaczące skutki dla budżetów i administracji krajowych?
Państwa członkowskie musiałyby wyznaczyć organy nadzorcze odpowiedzialne za wdrażanie wymogów legislacyjnych. Ich funkcja nadzorcza mogłaby opierać się na istniejących strukturach, na przykład obejmujących jednostki oceniające zgodność lub organy nadzoru rynku, ale wymagałaby odpowiedniej wiedzy specjalistycznej w zakresie technologii oraz zasobów. W zależności od istniejącej wcześniej struktury w każdym państwie członkowskim może to być od 1 do 25 ekwiwalentów pełnego czasu pracy na państwo członkowskie.
Czy wystąpią inne znaczące skutki?
Preferowany wariant znacznie ograniczyłby zagrożenia dla praw podstawowych obywateli oraz szerzej rozumianych wartości Unii, a ponadto zwiększy bezpieczeństwo niektórych produktów i usług z wbudowaną technologią sztucznej inteligencji lub samodzielnych zastosowań AI.
Proporcjonalność?
Wniosek jest proporcjonalny i niezbędny do osiągnięcia celów, ponieważ zastosowano w nim podejście oparte na analizie ryzyka, a obciążenia regulacyjne nakłada się tylko w przypadku, gdy systemy sztucznej inteligencji mogą stwarzać wysokie ryzyko dla praw podstawowych lub bezpieczeństwa. W przeciwnym razie nakłada się jedynie minimalne obowiązki w zakresie przejrzystości, w szczególności w zakresie dostarczania informacji w celu zasygnalizowania wykorzystania systemu sztucznej inteligencji przy kontaktowaniu się z ludźmi lub korzystania z technologii deepfake, jeśli nie stosuje się ich w uzasadnionych celach. Normy zharmonizowane oraz towarzyszące im wytyczne i narzędzia zapewniające przestrzeganie przepisów będą służyć wspieraniu dostawców i użytkowników w spełnieniu wymogów oraz zminimalizują ich koszty.
D. Działania następce
Kiedy nastąpi przegląd przyjętej polityki?
Komisja opublikuje sprawozdanie zawierające ocenę i przegląd ram pięć lat po dacie rozpoczęcia ich stosowania.