

Bruksela, 3 października 2022 r.
(OR. en)

Międzyinstytucjonalny numer
referencyjny:
2022/0303(COD)

13079/22
ADD 3

JUSTCIV 121
JAI 1260
CONSOM 243
COMPET 753
MI 703
FREMP 197
CODEC 1399
TELECOM 389
CYBER 311
DATAPROTECT 266

PISMO PRZEWODNIE

Od:	Sekretarz generalna Komisji Europejskiej (podpisała dyrektor Martine DEPREZ)
Data otrzymania:	29 września 2022 r.
Do:	Sekretariat Generalny Rady
Nr dok. Kom.:	SWD(2022) 320 final
Dotyczy:	DOKUMENT ROBOCZY SŁUŻB KOMISJI STRESZCZENIE SPRAWOZDANIA Z OCENY SKUTKÓW Towarzyszący dokumentowi: Wniosek DYREKTYWA PARLAMENTU EUROPEJSKIEGO I RADY w sprawie dostosowania przepisów dotyczących pozaumownej odpowiedzialności cywilnej do sztucznej inteligencji (dyrektywa w sprawie odpowiedzialności za sztuczną inteligencję)

Delegacje otrzymują w załączeniu dokument SWD(2022) 320 final.

Załącznik: SWD(2022) 320 final



KOMISJA
EUROPEJSKA

Bruksela, dnia 28.9.2022 r.
SWD(2022) 320 final

DOKUMENT ROBOCZY SŁUŻB KOMISJI
STRESZCZENIE SPRAWOZDANIA Z OCENY SKUTKÓW

[...]

Towarzyszący dokumentowi:

Wniosek
DYREKTYWA PARLAMENTU EUROPEJSKIEGO I RADY
w sprawie dostosowania przepisów dotyczących pozaumownej odpowiedzialności
cywilnej do sztucznej inteligencji
(dyrektywa w sprawie odpowiedzialności za sztuczną inteligencję)

{COM(2022) 496 final} - {SEC(2022) 344 final} - {SWD(2022) 318 final} -
{SWD(2022) 319 final}

Streszczenie oceny skutków
Ocena skutków inicjatywy dotyczącej odpowiedzialności cywilnej za szkody wyrządzone przez sztuczną inteligencję
A. Zasadność działań
Na czym polega problem i dlaczego jest to problem na szczeblu UE?
Upowszechnienie sztucznej inteligencji jest zarówno celem Komisji, jak i spodziewaną tendencją. Chociaż oczekuje się, że produkty/usługi oparte na sztucznej inteligencji będą bezpieczniejsze niż rozwiązana tradycyjne, wypadki nadal będą się zdarzać. Obowiązujące przepisy dotyczące odpowiedzialności, w szczególności przepisy krajowe oparte na zasadzie winy, nie są przystosowane do rozpatrywania roszczeń o odszkodowanie za szkody spowodowane przez produkty/usługi oparte na sztucznej inteligencji. Zgodnie z takimi przepisami poszkodowani muszą udowodnić niewłaściwe działanie/zaniechanie, którego dopuściła się osoba, która spowodowała szkodę. Cechy szczególne sztucznej inteligencji, w tym autonomia i brak przejrzystości (tzw. efekt czarnej skrzynki), sprawiają, że zidentyfikowanie osoby, która ponosi odpowiedzialność, oraz spełnienie wymogów uznania roszczenia z tytułu odpowiedzialności może być trudne lub nadmiernie kosztowne. Komisja chce uniknąć sytuacji, w której osoby poszkodowane z tytułu działania sztucznej inteligencji, tacy jak obywatele lub przedsiębiorstwa, byłiby chronieni w mniejszym stopniu niż poszkodowani z tytułu działania tradycyjnych technologii. Brak możliwości uzyskania rekompensaty z tytułu działania sztucznej inteligencji może wpłynąć na poziom zaufania do sztucznej inteligencji, a ostatecznie – na wykorzystanie produktów/usług na niej opartych. Nie jest pewne, w jaki sposób krajowe przepisy dotyczące odpowiedzialności mają być stosowane w kontekście swoistości sztucznej inteligencji. Ponadto, jeżeli rezultat postępowania byłby niesprawiedliwy dla poszkodowanego, sądy mogą stosować istniejące przepisy w sposób doraźny, aby osiągnąć sprawiedliwy rezultat. Spowoduje to brak pewności prawa. W efekcie przedsiębiorstwa będą miały trudności z przewidzeniem tego, w jaki sposób obowiązujące przepisy dotyczące odpowiedzialności będą stosowane w przypadku wystąpienia szkód. Doświadczą zatem trudności z oceną swojego narażenia na odpowiedzialność i zabezpieczeniem się przed takim narażeniem. Wpływ ten będzie jeszcze większy w przypadku przedsiębiorstw prowadzących działalność transgraniczną, ponieważ brak pewności będzie dotyczył różnych jurysdykcji. Wywrze to szczególny wpływ na MŚP, które nie mogą polegać na specjalistycznej wiedzy prawnej ani na rezerwach kapitałowych. Przewiduje się również, że jeśli UE nie podejmie żadnych działań, państwa członkowskie same dostosują swoje krajowe przepisy dotyczące odpowiedzialności do wyzwań związanych ze sztuczną inteligencją. Doprowadzi to do dalszego rozdrobnienia przepisów i wzrostu kosztów dla przedsiębiorstw prowadzących działalność transgraniczną.
Co należy osiągnąć?
Inicjatywa stanowi realizację priorytetu Komisji w zakresie transformacji cyfrowej. Jej nadrzędnym celem jest zatem promowanie wdrażania godnej zaufania sztucznej inteligencji, aby w pełni czerpać z niej korzyści. Dlatego też w białej księdze w sprawie sztucznej inteligencji za cel przyjęto stworzenie ekosystemu zaufania, aby przyczynić się do upowszechnienia sztucznej inteligencji. Inicjatywa dotycząca odpowiedzialności stanowi niezbędne następstwo przepisów bezpieczeństwa dostosowanych do sztucznej

inteligencji, a zatem uzupełnia akt w sprawie AI.

Inicjatywa w sprawie AI:

- zapewni, aby ochrona osób poszkodowanych w wyniku działania produktów/usług opartych na sztucznej inteligencji była równoważna ochronie osób poszkodowanych w wyniku działania technologii tradycyjnych;
- zwiększy pewność prawa w zakresie narażenia przedsiębiorstw opracowujących lub wykorzystujących sztuczną inteligencję na odpowiedzialność;
- służy zapobieganiu powstawania rozdrobnionych dostosowań krajowych przepisów dotyczących odpowiedzialności cywilnej w odniesieniu do sztucznej inteligencji.

Na czym polega wartość dodana podjęcia działań na poziomie UE (zasada pomocniczości)?

Wspieranie upowszechnienia sztucznej inteligencji w Europie oznacza konieczność otwarcia jednolitego rynku UE dla tych podmiotów gospodarczych, które chcą rozwijać sztuczną inteligencję lub wdrażać ją w swoich przedsiębiorstwach.

Warunkiem wstępnym jest zwiększenie pewności prawa i zapobieżenie rozdrobnieniu przepisów, do którego doszłoby, gdyby państwa członkowskie zaczęły samodzielnie dostosowywać przepisy krajowe w sposób rozbieżny.

Z ostrożnych szacunków wynika, że działanie na poziomie UE w kwestii odpowiedzialności za sztuczną inteligencję miałyby pozytywny wpływ w postaci 5–7 % wzrostu wartości produkcji w odpowiednim handlu transgranicznym w porównaniu ze scenariuszem bazowym.

B. Rozwiązania

Jakie są różne warianty działań służących osiągnięciu celów? Czy wskazano preferowany wariant? Jeżeli nie, dlaczego?

Wariant strategiczny nr 1: trzy środki mające na celu zmniejszenie ciężaru dowodu spoczywającego na poszkodowanych, którzy starają się udowodnić swoje roszczenie z tytułu odpowiedzialności cywilnej:

a) Harmonizacja sposobu ujawniania w postępowaniu sądowym informacji zarejestrowanych/udokumentowanych zgodnie z przepisami w zakresie bezpieczeństwa produktów zawartymi w akcie w sprawie AI, tak aby poszkodowany mógł wskazać i udowodnić, jakie działanie/zaniechanie doprowadziło do poniesionej szkody.

b) Jeżeli poszkodowany wykaże, że osoba ponosząca odpowiedzialność nie spełniła wymogów bezpieczeństwa określonych w akcie w sprawie AI mających na celu zapobieżenie szkodzie, możliwość zastosowania domniemania sądowego, zgodnie z którym niespełnienie tych wymogów stanowiło przyczynę powstania szkody. Osoba potencjalnie ponosząca odpowiedzialność miałaby możliwość wzruszenia takiego domniemania, np. poprzez udowodnienie, że do szkody doprowadziła inna przyczyna.

c) W przypadku, gdy jedynym sposobem, aby poszkodowany udowodnił swoje roszczenie z tytułu odpowiedzialności cywilnej, jest wykazanie, co zaszło wewnątrz systemu sztucznej inteligencji, takie obciążenie poszkodowanego zostałoby zmniejszone. Osoba potencjalnie ponosząca odpowiedzialność miałaby możliwość udowodnienia, że nie dopuściła się zaniedbania.

Wariant strategiczny nr 2: środki, o których mowa w wariantcie nr 1 + harmonizacja przepisów dotyczących odpowiedzialności na zasadzie ryzyka w odniesieniu do przypadków użycia sztucznej

inteligencji o szczególnym profilu ryzyka. Odpowiedzialność na zasadzie ryzyka oznacza, że osoba narażająca społeczeństwo na ryzyko (często ryzyko dla interesów prawnych o dużej wartości, takich jak życie, zdrowie, własność) i czerpiąca z tego korzyści ponosi odpowiedzialność w przypadku realizacji takiego ryzyka, np. odpowiedzialność właściciela samochodu. W takich przypadkach poszkodowany musi jedynie udowodnić, że zaistniała szkoda wynika z ryzyka stworzonego przez osobę ponoszącą odpowiedzialność. Wariant ten może zostać połączony z obowiązkowym ubezpieczeniem.

Wariant strategiczny nr 3: podejście etapowe (**preferowany wariant strategiczny**) obejmujące:

- etap pierwszy: środki, o których mowa w wariantcie nr 1;
- etap drugi: mechanizm przeglądu służący do ponownej oceny dotyczącej konieczności harmonizacji odpowiedzialności na zasadzie ryzyka w odniesieniu do przypadków użycia sztucznej inteligencji o szczególnym profilu ryzyka (potencjalnie w połączeniu z obowiązkowym ubezpieczeniem).

Jakie są opinie poszczególnych zainteresowanych stron? Jak kształtuje się poparcie dla poszczególnych wariantów?

Ogólnie rzecz biorąc, większość zainteresowanych stron przyznała, że zidentyfikowane problemy istnieją, i poparła działania na poziomie UE.

Obywatele UE, przedstawiciele organizacji konsumenckich i instytucji akademickich w przeważającej większości potwierdzili konieczność podjęcia przez UE działań mających na celu ograniczenie problemów związanych z ciężarem dowodu spoczywającym na osobach poszkodowanych. Chociaż przedstawiciele przedsiębiorstw uznali negatywne skutki braku pewności co do stosowania przepisów dotyczących odpowiedzialności, podeszli do sprawy bardziej zachowawczo i zwrócili się o ukierunkowaną interwencję, aby uniknąć ograniczania innowacji.

Podobnie było w przypadku wariantów strategicznych. Obywatele UE, przedstawiciele organizacji konsumenckich i instytucji akademickich zdecydowanie poparli, jako minimum, środki dotyczące ciężaru dowodu. Opowiedzieli się również za najsilniejszym środkiem, polegającym na harmonizacji odpowiedzialności na zasadzie ryzyka w połączeniu z obowiązkowym ubezpieczeniem.

Przedstawiciele przedsiębiorstw byli bardziej podzieleni, przy czym różnice zależały od wielkości przedsiębiorstwa. Odpowiedzialność na zasadzie ryzyka uznali za nieproporcjonalną. Harmonizacja w zakresie zmniejszenia ciężaru dowodu zyskała natomiast większe poparcie, w szczególności wśród MŚP. Przedsiębiorstwa przestrzegały jednak przed całkowitym przeniesieniem ciężaru dowodu.

C. Skutki wdrożenia preferowanego wariantu

Jakie korzyści przyniesie wdrożenie preferowanego wariantu lub – jeśli go nie wskazano – głównych wariantów?

Preferowany wariant strategiczny zapewniałby, aby osoby poszkodowane w wyniku stosowania produktów i usług opartych na sztucznej inteligencji (osoby fizyczne, przedsiębiorstwa i wszelkie inne podmioty publiczne lub prywatne) były chronione w nie mniejszym stopniu niż osoby poszkodowane w wyniku stosowania technologii tradycyjnych. Zwiększyłby on poziom zaufania do sztucznej inteligencji i do promowania jej upowszechniania.

Ponadto inicjatywa zwiększyłaby pewność prawa i zapobiegłaby rozdrobnieniu przepisów, co pomogłoby przedsiębiorstwom – w szczególności MŚP – chcącym w pełni wykorzystać potencjał jednolitego rynku

UE i wprowadzić w wymiarze transgranicznym produkty i usługi oparte na sztucznej inteligencji.

Inicjatywa poprawiłaby również warunki dla ubezpieczycieli umożliwiające im oferowanie objęcia ubezpieczeniem działań związanych ze sztuczną inteligencją, co ma zasadnicze znaczenie dla MŚP w zakresie zarządzania ryzykiem.

Jeśli chodzi o korzyści dla środowiska, oczekuje się, że inicjatywa przyczyni się do zwiększenia efektywności i pobudzenia innowacji w zakresie technologii przyjaznych dla środowiska.

Najnowocześniejsze produkty i usługi, których wspieranie jest celem inicjatywy, w większości nie są jeszcze dostępne na rynku. Proponowane środki wyprzedzają bieg wydarzeń, ponieważ dostosowują ramy prawne do szczególnych potrzeb i wyzwań związanych ze sztuczną inteligencją, tak aby stworzyć ekosystem oparty na zaufaniu i pewności prawa.

Ponieważ strategia ta ma charakter przyszłościowy, nie są dostępne dane wystarczające do ilościowego określenia skutków preferowanego wariantu polityki. Skutki oceniano zatem głównie pod względem jakościowym, przy uwzględnieniu wszystkich dostępnych danych, szacunków ekspertów i uwag zainteresowanych stron. W oparciu o racjonalne założenia podjęto pewne próby zastosowania podejścia ilościowego.

Szacuje się mianowicie, że preferowany wariant strategiczny doprowadziłby do zwiększenia wartości rynkowej sztucznej inteligencji w UE-27 do poziomu ok. 500 mln – 1,1 mld EUR w 2025 r. Ponadto z analizy mikroekonomicznej opartej na danych rynkowych dotyczących odkurzaczy automatycznych wynika, że inicjatywa spowodowałaby wzrost całkowitego dobrobytu o 30,11–53,74 mln EUR dla samej tej kategorii produktów w UE-27.

Jakie są koszty wdrożenia preferowanego wariantu lub – jeśli go nie wskazano – głównych wariantów?

Preferowany wariant polityki zapobiega lukom w zakresie odpowiedzialności spowodowanym cechami szczególnymi sztucznej inteligencji. Zapewniłby on, że w przypadkach, w których cechy szczególne sztucznej inteligencji uniemożliwiłyby poszkodowanemu udowodnienie niezbędnych faktów, zamiast poszkodowanego koszty ponosiłaby osoba ponosząca odpowiedzialność za spowodowanie szkody.

Jest to zgodne z jednym z podstawowych celów prawa dotyczącego odpowiedzialności, tj. zapewnieniem, aby osoba, która wyrządza szkodę innej osobie w sposób niezgodny z prawem, zrekompensowała poszkodowanemu wyrządzoną szkodę. Stałym celem polityki Komisji jest również zapewnienie, aby poziom ochrony osób, które poniosły szkody spowodowane przez systemy sztucznej inteligencji był równoważny poziomowi ochrony osób, które poniosły szkody spowodowane przez technologie tradycyjne. Prowadzi to do bardziej efektywnej alokacji kosztów, aby ponosiła je osoba, która rzeczywiście spowodowała szkodę i jest w stanie najlepiej zapobiec jej wystąpieniu.

Osoby potencjalnie ponoszące odpowiedzialność (w szczególności przedsiębiorcy działający na rynku sztucznej inteligencji) prawdopodobnie zostaną objęte ubezpieczeniem. Rozwiązania ubezpieczeniowe pozwalają rozłożyć ciężar odpowiedzialności na wszystkich ubezpieczonych, a tym samym ograniczyć koszty osób ponoszących odpowiedzialność do wysokości rocznych składek ubezpieczeniowych. Koszt wypłacenia osobie poszkodowanej odszkodowania wyrażałby się zatem dla ubezpieczonych osób ponoszących odpowiedzialność jedynie jako marginalny wzrost ich składek ubezpieczeniowych.

Wysoce solidne i precyzyjne określenie kosztów nie było możliwe, ponieważ najnowocześniejsze produkty i usługi promowane w ramach tej inicjatywy w większości nie są jeszcze dostępne na rynku.

W oparciu o dostępne dane, analizę ekspertów i uzasadnione założenia oszacowano, że preferowany wariant polityki może doprowadzić do zwiększenia rocznej całkowitej kwoty składek na ubezpieczenie od ogólnej odpowiedzialności cywilnej w UE o 5,35 mln EUR do 16,1 mln EUR.

Jakie są skutki dla MŚP i konkurencyjności?

Dzięki poprawie warunków funkcjonowania rynku wewnętrznego produktów i usług opartych na sztucznej inteligencji inicjatywa miałaby pozytywny wpływ na konkurencyjność przedsiębiorstw działających na europejskim rynku sztucznej inteligencji. Przedsiębiorstwa te stałyby się bardziej konkurencyjne w skali globalnej, co wzmocni pozycję UE w stosunku do jej konkurentów w światowym wyścigu w dziedzinie sztucznej inteligencji (przede wszystkim USA i Chin). Ponieważ sztuczna inteligencja jest przekrojową technologią wspomagającą, korzyści te nie byłyby ograniczone jedynie do pewnych konkretnych sektorów, ale miałyby zastosowanie – choć w różnym stopniu – we wszystkich sektorach, w których rozwija się lub wykorzystuje sztuczną inteligencję.

Zwiększenie pewności prawa i ograniczenie rozdrobnienia przepisów stanowiłoby większą korzyść dla MŚP niż dla innych zainteresowanych stron, ponieważ to właśnie MŚP są w większym stopniu dotknięte tymi problemami w ramach obecnych przepisów dotyczących odpowiedzialności. Inicjatywa ta poprawiłaby warunki w szczególności dla tych MŚP, które chciałyby wprowadzić na rynek w innych państwach członkowskich produkty lub usługi oparte na sztucznej inteligencji. Ma to kluczowe znaczenie, ponieważ unijny rynek sztucznej inteligencji jest w dużej mierze napędzany przez MŚP, które opracowują, wdrażają lub wykorzystują technologie sztucznej inteligencji.

MŚP skorzystałyby z tej inicjatywy również w charakterze osób poszkodowanych wynikiem działania sztucznej inteligencji, ponieważ mogłyby liczyć na zmniejszenie ciężaru dowodu przy ubieganiu się o odszkodowanie.

Czy przewiduje się znaczące skutki dla budżetów i administracji krajowych?

Nie przewiduje się znaczących skutków dla budżetów i administracji krajowych.

Przewidywane środki zmniejszające ciężar dowodu spoczywający na poszkodowanych można by bez przeszkód włączyć do istniejących ram proceduralnych i ram odpowiedzialności cywilnej państw członkowskich.

Państwa członkowskie będą musiały składać sprawozdania z realizacji inicjatywy i dostarczać określonych informacji na potrzeby ukierunkowanego przeglądu dokonywanego przez Komisję. Te wymogi w zakresie sprawozdawczości będą jednak ograniczone do informacji dostępnych w istniejących już bazach danych państw członkowskich oraz do informacji zgłaszanych na podstawie innych instrumentów prawnych (np. aktu w sprawie AI lub dyrektywy w sprawie ubezpieczeń pojazdów), co pozwoli osiągnąć synergii i zapewnić spójność przyszłych środków politycznych w różnych obszarach.

Czy wystąpią inne znaczące skutki?

Prawa podstawowe: Inicjatywa przyczyni się do wspierania skutecznego egzekwowania praw podstawowych na drodze prywatnoprawnej oraz do zachowania prawa do skutecznego środka odwoławczego w przypadku urzeczywistnienia się zagrożeń dla praw podstawowych związanych ze sztuczną inteligencją (np. dyskryminacji).

Wymiar międzynarodowy: Przedstawienie wyważonego podejścia dotyczącego odpowiedzialności za szkody wyrządzone przez sztuczną inteligencję, stwarza dla UE możliwość opracowania globalnego

punktu odniesienia i przedstawiania swojego podejścia jako rozwiązania globalnego, co w ostatecznym rozrachunku przyniosłoby przewagę konkurencyjną „AI made in Europe”.

Proporcjonalność?

Preferowany wariant ma na celu przygotowanie podstaw dla rozwoju i wykorzystania sztucznej inteligencji, przy jednoczesnym osiągnięciu głównego celu, jakim jest promowanie jej upowszechnienia w UE.

Wariant ten nie będzie jednak wykraczał poza to, co konieczne. Po pierwsze, interwencja UE jest ukierunkowana, ponieważ zmniejszy ona jedynie ciężar dowodu spoczywający na poszkodowanych. Zharmonizuje ona jedynie te kwestie dotyczące odpowiedzialności, które zmieniają się w wyniku powstania sztucznej inteligencji, a inne elementy, takie jak określenie winy i związku przyczynowego, nadal będą podlegały istniejącym przepisom krajowym.

Po drugie, preferowany wariant odkłada ocenę potrzeby harmonizacji odpowiedzialności na zasadzie ryzyka na późniejszy etap, kiedy będzie można zebrać więcej informacji na temat sztucznej inteligencji i jej zastosowań (zob. poniżej).

Po trzecie, w ramach preferowanego wariantu przyjęta zostanie zasada minimalnej harmonizacji. Minimalna harmonizacja nie zapewnia wprawdzie całkowicie równych warunków działania, ale gwarantuje, że nowe przepisy będzie można bez przeszkód włączyć do istniejących ram prawnych odpowiedzialności cywilnej w każdym państwie członkowskim.

Dzięki temu państwa członkowskie będą mogły włączyć do swojego prawa krajowego ukierunkowane interwencje unijne dokonane w ramach wariantu preferowanego, inicjatywa ta zwiększy w całej UE pewność prawa, zapobiegnie dalszemu rozdrobnieniu przepisów i zapewni skuteczną ochronę poszkodowanych, porównywalną z poziomem ochrony w przypadku innych rodzajów szkód.

D. Działania następne

Kiedy nastąpi przegląd przyjętej polityki?

Na preferowany wariant strategiczny składa się podejście etapowe: najpierw wprowadzone zostaną środki zmniejszające ciężar dowodu spoczywający na poszkodowanym, a następnie, na podstawie klauzuli przeglądowej, po pięciu latach dokonana zostanie ocena sytuacji. Proces ten pozwoli Komisji ocenić w świetle rozwoju technologii i jej zastosowań, czy oprócz środków zmniejszających ciężar dowodu potrzebna jest również harmonizacja odpowiedzialności na zasadzie ryzyka oraz wprowadzenie obowiązkowego ubezpieczenia.