

Bruksela, 24 lutego 2020 r.
(OR. en)

6266/20

JAI 149
COPEN 60
CYBER 27
DATAPROTECT 23
EJUSTICE 22
COSI 28
IXIM 30
ENFOPOL 55
FREMP 15
TELECOM 24
RELEX 151
MI 49
COMPET 57

PISMO PRZEWODNIE

Od:	Sekretarz Generalny Komisji Europejskiej, podpisał dyrektor Jordi AYET PUIGARNAU
Data otrzymania:	20 lutego 2020 r.
Do:	Jeppe TRANHOLM-MIKKELSEN, Sekretarz Generalny Rady Unii Europejskiej
Nr dok. Kom.:	COM(2020) 65 final
Dotyczy:	BIAŁA KSIĘGA w sprawie sztucznej inteligencji Europejskie podejście do doskonałości i zaufania

Delegacje otrzymują w załączeniu dokument COM(2020) 65 final.

Zał.: COM(2020) 65 final



KOMISJA
EUROPEJSKA

Bruksela, dnia 19.2.2020 r.
COM(2020) 65 final

BIAŁA KSIĘGA

w sprawie sztucznej inteligencji
Europejskie podejście do doskonałości i zaufania

Biała księga w sprawie sztucznej inteligencji

Europejskie podejście do doskonałości i zaufania

Sztuczna inteligencja (AI) rozwija się w szybkim tempie. Zmieni nasze życie dzięki poprawie opieki zdrowotnej (np. bardziej precyzyjna diagnostyka, lepsze zapobieganie chorobom), zwiększeniu wydajności rolnictwa, przyczynieniu się do adaptacji do zmiany klimatu i jej łagodzenia, poprawie wydajności systemów produkcji w wyniku konserwacji predykcyjnej, zwiększeniu bezpieczeństwa Europejczyków oraz na wiele innych sposobów, których nie jesteśmy sobie w stanie nawet wyobrazić. Jednocześnie sztuczna inteligencja wiąże się z szeregiem potencjalnych zagrożeń takich jak nieprzejrzyste podejmowanie decyzji, dyskryminacja ze względu na płeć lub inne czynniki, ingerencja w nasze życie prywatne czy wykorzystanie w celach przestępczych.

W obliczu intensywnej konkurencji na świecie potrzebne jest solidne europejskie podejście oparte na europejskiej strategii na rzecz sztucznej inteligencji przedstawionej w kwietniu 2018 r.¹ Aby osiągnąć korzyści i stawić czoła wyzwaniom, jakie niesie ze sobą sztuczna inteligencja, UE musi działać jednomyślnie oraz opracować własny sposób działania w celu promowania rozwoju i wprowadzania sztucznej inteligencji w oparciu o wartości europejskie.

Komisja zobowiązała się do umożliwienia przełomowych badań naukowych, utrzymania wiodącej pozycji UE w zakresie technologii oraz zapewnienia, aby nowe technologie służyły wszystkim Europejczykom i przyczyniały się do poprawy ich życia przy jednoczesnym poszanowaniu ich praw.

Przewodnicząca Komisji Ursula von der Leyen ogłosiła w swoich wytycznych politycznych² skoordynowane europejskie podejście do społecznych i etycznych konsekwencji sztucznej inteligencji, a także podjęcie refleksji nad lepszym wykorzystaniem dużych zbiorów danych na rzecz innowacji.

W związku z tym Komisja popiera podejście regulacyjne i inwestycyjne, którego podwójnym celem jest promowanie stosowania sztucznej inteligencji i zajęcie się zagrożeniami związanymi z niektórymi zastosowaniami tej nowej technologii. Celem niniejszej białej księgi jest określenie wariantów strategicznych dotyczących sposobów osiągnięcia tych założeń. Nie zajęto się w niej rozwojem i wykorzystaniem sztucznej inteligencji do celów wojskowych. Komisja zachęca państwa członkowskie, inne instytucje europejskie i wszystkie zainteresowane strony, w tym przemysł, partnerów społecznych, organizacje społeczeństwa obywatelskiego, naukowców, ogół społeczeństwa i każdą zainteresowaną osobę, do wyrażenia opinii na temat poniższych wariantów strategicznych i wniesienia wkładu w przyszłe decyzje Komisji w tej dziedzinie.

1. WPROWADZENIE

Technologia cyfrowa staje się coraz bardziej centralną częścią każdego aspektu życia, dlatego ludzie powinni móc jej ufać. Wiarygodność tej technologii jest warunkiem wstępnym jej upowszechnienia. Jest to szansa dla Europy ze względu na jej silne przywiązanie do wartości i praworządności oraz potwierdzoną zdolność tworzenia bezpiecznych, niezawodnych i wyrafinowanych produktów i usług – od aeronautyki po energetykę, motoryzację i sprzęt medyczny.

Obecny i przyszły zrównoważony wzrost gospodarczy i dobrobyt społeczny Europy w coraz większym stopniu opierają się na wartości wytworzonej przez dane. Sztuczna inteligencja to jedno

¹ Sztuczna inteligencja dla Europy, COM(2018) 237 final.

² https://ec.europa.eu/commission/sites/beta-political/files/political-guidelines-next-commission_pl.pdf.

z najważniejszych zastosowań gospodarki opartej na danych. Obecnie większość danych jest związana z konsumentami i są one przechowywane i przetwarzane w centralnej infrastrukturze opartej na chmurze obliczeniowej. Natomiast znaczna część przyszłych, znacznie bogatszych danych pochodzić będzie z przemysłu, przedsiębiorstw i sektora publicznego i będzie przechowywana w różnych systemach, w szczególności w urządzeniach obliczeniowych pracujących na obrzeżach sieci. Otwiera to nowe możliwości dla Europy, która ma silną pozycję w branży cyfrowej i w zastosowaniach między przedsiębiorstwami (business-to-business), ale ma stosunkowo słabą pozycję, jeśli chodzi o platformy konsumenckie.

Sztuczna inteligencja to zbiór technologii łączących dane, algorytmy i moc obliczeniową. Główną siłą napędową obecnego rozwoju sztucznej inteligencji są postępy w dziedzinie obliczeń i coraz większa dostępność danych. Europa może połączyć swoją przewagę technologiczną i przemysłową z wysokiej jakości infrastrukturą cyfrową i ramami regulacyjnymi opartymi na europejskich wartościach podstawowych, aby **stać się światowym liderem w dziedzinie innowacji w gospodarce opartej na danych i jej zastosowaniach**, jak określono w europejskiej strategii w zakresie danych³. Na tej podstawie można rozwijać ekosystem sztucznej inteligencji zapewniający korzyści płynące z tej technologii dla całego społeczeństwa i gospodarki:

- dla **obywateli** – na przykład lepsza opieka zdrowotna, rzadziej psujący się sprzęt AGD, bezpieczniejsze i czystsze systemy transportu, lepsze usługi publiczne;
- dla rozwoju **przedsiębiorstw** – na przykład nowa generacja produktów i usług w obszarach, w których Europa jest szczególnie silna (sektor maszyn, transportu, cyberbezpieczeństwa, rolnictwa, zielona gospodarka o obiegu zamkniętym, sektor opieki zdrowotnej i sektory o wysokiej wartości dodanej, takie jak moda i turystyka); oraz
- dla usług **interesu publicznego** – na przykład zmniejszenie kosztów świadczenia usług (transport, edukacja, energia i gospodarowanie odpadami), poprawa zrównoważonego charakteru produktów⁴ oraz wyposażenie organów egzekwowania prawa w adekwatne narzędzia zapewniające bezpieczeństwo obywateli⁵ przy zachowaniu odpowiednich zabezpieczeń w odniesieniu do ich praw i swobód.

Biorąc pod uwagę znaczny wpływ, jaki sztuczna inteligencja może wywierać na nasze społeczeństwo oraz konieczność zbudowania zaufania do niej, bardzo ważne jest, aby europejska AI opierała się na naszych wartościach i prawach podstawowych, takich jak godność ludzka i ochrona prywatności.

Ponadto należy brać pod uwagę wpływ systemów AI nie tylko z perspektywy indywidualnej, lecz również z perspektywy społeczeństwa jako całości. Systemy sztucznej inteligencji mogą odegrać znaczącą rolę w osiągnięciu celów zrównoważonego rozwoju oraz we wspieraniu procesu demokratycznego i praw socjalnych. Dzięki swoim niedawnym wnioskom w sprawie Europejskiego Zielonego Ładu⁶ Europa przoduje w stawianiu czoła wyzwaniom związanym z klimatem

³ COM(2020) 66 final.

⁴ Sztuczna inteligencja i cyfryzacja ogółem są kluczowymi czynnikami umożliwiającymi osiągnięcie ambitnych celów Europejskiego Zielonego Ładu. Jednocześnie szacuje się, że obecny ślad środowiskowy sektora ICT wynosi ponad 2 % wszystkich emisji na świecie. W europejskiej strategii cyfrowej towarzyszącej niniejszej białej księdze zaproponowano środki na rzecz ekologicznej transformacji sektora cyfrowego.

⁵ Narzędzia w zakresie sztucznej inteligencji mogą przyczynić się do lepszej ochrony obywateli UE przed przestępczością i aktami terrorystycznymi. Takie narzędzia mogłyby na przykład pomóc w identyfikowaniu internetowej propagandy terrorystycznej, wykrywaniu podejrzanych transakcji sprzedaży niebezpiecznych produktów, identyfikowaniu niebezpiecznych ukrytych przedmiotów lub nielegalnych substancji lub produktów, czy też oferować pomoc obywatelom w sytuacjach nadzwyczajnych i wspomagać służby interwencyjne.

⁶ COM(2019) 640 final.

i środowiskiem. Technologie cyfrowe, takie jak sztuczna inteligencja, są kluczowym czynnikiem umożliwiającym osiągnięcie celów Zielonego Ładu. Biorąc pod uwagę rosnące znaczenie sztucznej inteligencji należy odpowiednio uwzględnić wpływ systemów AI na środowisko w całym cyklu życia i w całym łańcuchu dostaw, np. w odniesieniu do wykorzystania zasobów do szkolenia algorytmów i przechowywania danych.

Wspólne europejskie podejście do sztucznej inteligencji jest konieczne do osiągnięcia wystarczającej skali i uniknięcia rozdrobnienia jednolitego rynku. Wprowadzenie inicjatyw krajowych może zagrozić pewności prawa, osłabić zaufanie obywateli i uniemożliwić powstanie dynamicznego przemysłu europejskiego.

W niniejszej białej księdze przedstawiono warianty strategiczne umożliwiające bezpieczny rozwój godnej zaufania sztucznej inteligencji w Europie, przy pełnym poszanowaniu wartości i praw obywateli europejskich. Główne elementy niniejszej białej księgi to:

- Ramy polityczne określające środki służące połączeniu wysiłków na szczeblu europejskim, krajowym i regionalnym. Dzięki partnerstwu sektora publicznego i prywatnego ramy te powinny zmobilizować zasoby w celu osiągnięcia „**ekosystemu doskonałości**” wzdłuż całego łańcucha wartości, poczynając od badań naukowych i innowacji, a także stworzyć odpowiednie zachęty do przyspieszenia przyjmowania rozwiązań opartych na sztucznej inteligencji, w tym przez małe i średnie przedsiębiorstwa (MŚP).
- Kluczowe elementy przyszłych ram regulacyjnych dotyczących sztucznej inteligencji w Europie, które stworzą wyjątkowy „**ekosystem zaufania**”. W tym celu muszą one zapewniać poszanowanie przepisów UE, w tym przepisów służących ochronie praw podstawowych i praw konsumentów, w szczególności w odniesieniu do wykorzystywanych w UE systemów AI charakteryzujących się wysokim ryzykiem⁷. Budowanie ekosystemu zaufania jest celem politycznym samym w sobie i powinno zachęcać obywateli do stosowania sztucznej inteligencji oraz oferować przedsiębiorstwom i organizacjom publicznym pewność prawa umożliwiającą innowacyjność z wykorzystaniem AI. Komisja zdecydowanie popiera podejście, w którego centrum jest człowiek i które będzie opierać się na komunikacji w sprawie budowania zaufania do sztucznej inteligencji ukierunkowanej na człowieka⁸. Komisja uwzględni również wkład uzyskany w fazie pilotażowej prac nad wytycznymi dotyczącymi etyki przygotowanymi przez grupę ekspertów wysokiego szczebla ds. sztucznej inteligencji.

Europejska strategia w zakresie danych, która towarzyszy niniejszej białej księdze, ma na celu umożliwienie Europie stania się najbardziej atrakcyjną, bezpieczną i dynamiczną gospodarką sprawnie wykorzystującą dane – wyposażając Europę w dane umożliwiające podejmowanie lepszych decyzji i poprawę życia wszystkich obywateli. W strategii określono szereg środków z zakresu polityki, w tym mobilizację inwestycji prywatnych i publicznych niezbędnych do osiągnięcia tego celu. Wpływ sztucznej inteligencji, internetu rzeczy i innych technologii cyfrowych na przepisy dotyczące bezpieczeństwa i odpowiedzialności został przeanalizowany w sprawozdaniu Komisji towarzyszącym niniejszej białej księdze.

⁷ Mimo że konieczne mogą być dalsze ustalenia w celu zapobiegania i zwalczania stosowania AI w celach przestępczych, wykraczają one poza zakres niniejszej białej księgi.

⁸ COM(2019) 168.

2. WYKORZYSTANIE MOCNYCH STRON RYNKÓW PRZEMYSŁOWYCH I PROFESJONALNYCH

Europa jest dobrze przygotowana do korzystania z potencjału sztucznej inteligencji nie tylko jako użytkownik, ale również jako twórca i producent tej technologii. Znajdują się w niej doskonałe ośrodki badawcze i innowacyjne start-upy. Odgrywa ona również wiodącą rolę w dziedzinie robotyki i konkurencyjnych sektorów produkcji i usług, od motoryzacji po opiekę zdrowotną, energetykę, usługi finansowe czy rolnictwo. Europa rozwinęła silną infrastrukturę obliczeniową (np. komputery na potrzeby obliczeń wielkiej skali), co ma zasadnicze znaczenie dla funkcjonowania sztucznej inteligencji. Posiada również duże ilości danych publicznych i przemysłowych, których potencjał nie jest obecnie w pełni wykorzystywany i dysponuje uznanymi zasobami przemysłowymi w zakresie bezpiecznych systemów cyfrowych o niskim poborze mocy, które są niezbędne do dalszego rozwoju sztucznej inteligencji.

Wykorzystanie zdolności UE do inwestowania w technologie i infrastrukturę nowej generacji oraz w kompetencje cyfrowe, takie jak umiejętność korzystania z danych, zwiększy technologiczną suwerenność Europy w zakresie kluczowych technologii wspomagających i infrastruktury gospodarki opartej na danych. Infrastruktura powinna wspierać tworzenie europejskich puli danych umożliwiających powstanie godnej zaufania sztucznej inteligencji, np. opartej na wartościach i zasadach europejskich.

Europa powinna wykorzystać swoje silne strony, aby umocnić pozycję w ekosystemach i wzdłuż łańcucha wartości, począwszy od niektórych sektorów produkcji sprzętu, poprzez oprogramowanie, aż po usługi. W pewnym stopniu już się to dzieje. W Europie wytwarza się ponad jedną czwartą wszystkich usługowych robotów przemysłowych i profesjonalnych (wykorzystywanych np. w sektorach rolnictwa precyzyjnego, bezpieczeństwa, zdrowia, logistyki). Odgrywa ona ważną rolę w opracowywaniu i wykorzystywaniu oprogramowania dla przedsiębiorstw i organizacji (np. oprogramowanie wykorzystywane między przedsiębiorstwami służące planowaniu zasobów przedsiębiorstwa czy oprogramowanie do celów projektowania i zastosowań konstrukcyjnych), jak również oprogramowania do celów administracji elektronicznej i „inteligentnych przedsiębiorstw”.

Europa przoduje we wdrażaniu sztucznej inteligencji w przemyśle wytwórczym. Ponad połowa jej czołowych producentów stosuje sztuczną inteligencję przynajmniej w jednym aspekcie procesu produkcji⁹.

Jednym z powodów silnej pozycji Europy w dziedzinie badań naukowych jest unijny program finansowania, który okazał się niezwykle istotny, jeśli chodzi o łączenie działań, unikanie ich powielania oraz pozyskiwanie inwestycji publicznych i prywatnych w państwach członkowskich. W ciągu ostatnich trzech lat unijne środki finansowe przeznaczone na badania naukowe i innowacje w dziedzinie sztucznej inteligencji wzrosły do 1,5 mld EUR, co oznacza wzrost o 70 % w porównaniu z poprzednim okresem.

Inwestycje w badania i innowacje w Europie stanowią jednak nadal niewielką część inwestycji publicznych i prywatnych w porównaniu z innymi regionami świata. W 2016 r. w Europie w sztuczną inteligencję zainwestowano około 3,2 mld EUR, w porównaniu z około 12,1 mld EUR w Ameryce Północnej i 6,5 mld EUR w Azji¹⁰. Europa musi zatem znacznie zwiększyć poziom inwestycji. Skoordynowany plan w sprawie sztucznej inteligencji¹¹ opracowany wraz z państwami członkowskimi

⁹ Kolejne są Japonia (30 %) i USA (28 %). Źródło: CapGemini (2019).

¹⁰ 10 imperatives for Europe in the age of AI and automation (10 imperatywów dla Europy w erze AI i automatyzacji), McKinsey (2017).

¹¹ COM(2018) 795.

okazuje się dobrym punktem wyjścia do budowania ściślejszej współpracy w zakresie sztucznej inteligencji w Europie oraz tworzenia synergii w celu maksymalizacji inwestycji w łańcuch wartości sztucznej inteligencji.

3. WYKORZYSTANIE PRZYSZŁYCH MOŻLIWOŚCI: KOLEJNA FAŁA DANYCH

Europa znajduje się obecnie na słabszej pozycji w dziedzinie zastosowań konsumenckich i platform internetowych, co prowadzi do niekorzystnej sytuacji konkurencyjnej jeśli chodzi o dostęp do danych. Nadchodzą jednak duże zmiany w odniesieniu do wartości i ponownego wykorzystania danych w różnych sektorach. Ilość danych generowanych na świecie gwałtownie rośnie – z 33 zettabajtów w 2018 r. do prognozowanych 175 zettabajtów w 2025 r.¹² Nowe fale danych stwarzają za każdym razem Europie szansę ugruntowania pozycji w gospodarce sprawnie wykorzystującej dane oraz stania się światowym liderem w tej dziedzinie. Ponadto sposób przechowywania i przetwarzania danych diametralnie się zmieni w ciągu najbliższych pięciu lat. Obecnie 80 % operacji przetwarzania i analizy danych, które mają miejsce w chmurze, odbywa się w centrach danych i w ramach scentralizowanej infrastruktury obliczeniowej, a 20 % w inteligentnych przedmiotach podłączonych do internetu, takich jak samochody, sprzęty gospodarstwa domowego czy roboty produkcyjne, oraz w obiektach przetwarzania danych znajdujących się w pobliżu użytkownika (tzw. *edge computing*). Oczekuje się, że do 2025 r. proporcja ta ulegnie zasadniczej zmianie¹³.

Europa jest światowym liderem w dziedzinie elektroniki o niskim poborze mocy, która ma kluczowe znaczenie dla kolejnej generacji wyspecjalizowanych procesorów na potrzeby AI. Rynek ten jest obecnie zdominowany przez podmioty spoza UE. Mogłoby to się zmienić dzięki inicjatywom takim jak europejska inicjatywa dotycząca procesorów, która odnosi się do rozwoju systemów obliczeniowych o niskim poborze mocy zarówno na potrzeby przetwarzania danych na obrzeżach sieci, jak i obliczeń wielkiej skali nowej generacji, a także dzięki działalności wspólnego przedsięwzięcia w obszarze kluczowych technologii cyfrowych, która rozpocznie się w 2021 r. Ponadto Europa przoduje w rozwiązaniach neuromorficznych¹⁴, które doskonale nadają się do automatyzacji procesów przemysłowych (przemysł 4.0) i transportu. Mogą one zwiększyć efektywność energetyczną o kilka rzędów wielkości.

Najnowsze osiągnięcia w dziedzinie obliczeń kwantowych pozwolą na gwałtowny wzrost zdolności obliczeniowych¹⁵. Europa może stać się liderem tej technologii dzięki swoim osiągnięciom akademickim dotyczącym kwantowych technologii obliczeniowych, a także dzięki silnej pozycji europejskiego przemysłu w dziedzinie symulatorów kwantowych i środowisk programowania obliczeń kwantowych. Europejskie inicjatywy mające na celu zwiększenie dostępności obiektów przeprowadzających testy i doświadczenia z zastosowaniem kwantowych technologii obliczeniowych pomogą w stosowaniu nowych rozwiązań kwantowych w różnych sektorach przemysłu i środowiska akademickiego.

Jednocześnie Europa będzie nadal przewodzić postępom w zakresie algorytmicznych podstaw sztucznej inteligencji, opierając się na swojej własnej doskonałości naukowej. Istnieje potrzeba budowania połączeń między dziedzinami, które obecnie funkcjonują osobno, takimi jak uczenie się maszyn i uczenie głębokie (które charakteryzują się ograniczoną interoperacyjnością i potrzebą

¹² IDC (2019).

¹³ Gartner (2017).

¹⁴ Rozwiązania neuromorficzne oznaczają każdy bardzo duży system układów scalonych imitujących architekturę biologiczną występującą w układzie nerwowym.

¹⁵ Komputery kwantowe będą mogły przetwarzać – w czasie krótszym niż sekundy – wielokrotnie większe zbiory danych niż najszybsze dziś komputery, co umożliwi rozwój nowych zastosowań AI we wszystkich sektorach.

dysponowania dużą ilością danych w celu szkolenia modeli i uczenia się poprzez korelacje) oraz podejściami symbolicznymi (gdzie zasady są tworzone w drodze ludzkiej interwencji). Połączenie rozumowania symbolicznego z głębokimi sieciami neuronowymi może przyczynić się do poprawy wytłumaczalności wyników AI.

4. EKOSYSTEM DOSKONAŁOŚCI

Aby zbudować ekosystem doskonałości, który może wspierać rozwój i stosowanie sztucznej inteligencji w całej gospodarce UE i w administracji publicznej, należy zintensyfikować działania na wielu poziomach.

A. WSPÓŁPRACA Z PAŃSTWAMI CZŁONKOWSKIMI

Realizując strategię na rzecz sztucznej inteligencji przyjętą w kwietniu 2018 r.¹⁶, w grudniu 2018 r. Komisja przedstawiła opracowany wspólnie z państwami członkowskimi skoordynowany plan wspierania rozwoju i stosowania sztucznej inteligencji w Europie¹⁷.

W planie zaproponowano ok. 70 wspólnych działań na rzecz ściślejszej i skuteczniejszej współpracy między państwami członkowskimi a Komisją w kluczowych obszarach, takich jak badania, inwestycje, wprowadzanie na rynek, umiejętności i talent, dane i współpraca międzynarodowa. Oczekuje się, że plan będzie realizowany do 2027 r. oraz będzie regularnie monitorowany i poddawany przeglądom.

Celem jest zmaksymalizowanie wpływu inwestycji w dziedzinie badań naukowych, innowacji i wdrażania, ocena krajowych strategii w zakresie AI oraz rozwijanie skoordynowanego planu w sprawie sztucznej inteligencji wraz z państwami członkowskimi:

- *Działanie nr 1: Uwzględniając wyniki publicznych konsultacji w sprawie białej książki, Komisja zaproponuje państwom członkowskim przeprowadzenie do końca 2020 r. przeglądu skoordynowanego planu*

Unijne finansowanie w dziedzinie sztucznej inteligencji powinno przyciągać i scalać inwestycje w obszarach, w których wymagane działania wykraczają poza to, co może osiągnąć pojedyncze państwo członkowskie. Celem jest przyciągnięcie w następnym dziesięcioleciu w całej UE rocznie ponad 20 mld EUR¹⁸ łącznych inwestycji w sztuczną inteligencję. Mając na uwadze potrzeby słabiej rozwiniętych regionów i obszarów wiejskich oraz w celu zmobilizowania inwestycji prywatnych i publicznych UE udostępni środki z programów „Cyfrowa Europa” i „Horyzont Europa” oraz europejskich funduszy strukturalnych i inwestycyjnych.

W skoordynowanym planie można również uwzględnić kwestię dobrostanu społecznego i środowiskowego jako kluczową zasadę w odniesieniu do sztucznej inteligencji. Systemy sztucznej inteligencji mają pomóc w rozwiązaniu najpilniejszych problemów, takich jak zmiana klimatu i degradacja środowiska. Musi się to odbywać w sposób przyjazny dla środowiska. Sztuczna inteligencja może i powinna krytycznie oceniać zużycie zasobów i energii oraz być przeszkolona w zakresie dokonywania wyborów, które są przyjazne dla środowiska. Komisja we współpracy z państwami członkowskimi rozważy możliwości wspierania i promowania rozwiązań w zakresie AI służących osiągnięciu tych celów.

¹⁶ Sztuczna inteligencja dla Europy, COM(2018) 237.

¹⁷ Skoordynowany plan w sprawie sztucznej inteligencji, COM (2018) 795.

¹⁸ COM(2018) 237.

B. UKIERUNKOWANIE DZIAŁAŃ SPOŁECZNOŚCI BADAWCZEJ I INNOWACYJNEJ

Europa nie może sobie pozwolić na utrzymanie obecnego rozproszonego krajobrazu ośrodków kompetencji w sytuacji, gdy żaden z nich nie osiągnie skali niezbędnej do konkurowania z czołowymi instytutami na świecie. Konieczne jest stworzenie większej synergii i sieci między różnymi europejskimi ośrodkami badawczymi zajmującymi się sztuczną inteligencją oraz skoordynowanie ich działań w celu poprawy doskonałości, zatrzymywania i przyciągania najlepszych naukowców oraz opracowywania najlepszych technologii. Europa potrzebuje sztandarowego centrum badań naukowych, innowacji i wiedzy fachowej, które będzie koordynować te działania, stanowić światowy wzór doskonałości w dziedzinie sztucznej inteligencji oraz przyciągać inwestycje i największe talenty w tej dziedzinie.

Centra i sieci powinny koncentrować się w sektorach, w których Europa może stać się światowym liderem, takich jak przemysł, zdrowie, transport, finanse, rolno-spożywcze łańcuchy wartości, energia/środowisko, leśnictwo, obserwacja Ziemi i przestrzeń kosmiczna. We wszystkich tych dziedzinach toczy się wyścig o pozycję światowego lidera, a Europa posiada znaczny potencjał, wiedzę i doświadczenie¹⁹. Równie ważne jest tworzenie ośrodków badawczych i doświadczalnych w celu wspierania rozwoju i późniejszego wprowadzania nowych zastosowań w dziedzinie sztucznej inteligencji.

- *Działanie nr 2: Komisja ułatwi tworzenie centrów doskonałości i centrów badawczych, które mogą łączyć inwestycje europejskie, krajowe i prywatne, być może przy użyciu nowego instrumentu prawnego. Jako część wieloletnich ram finansowych na lata 2021–2027 Komisja zaproponowała przeznaczenie ambitnej kwoty na wsparcie światowych ośrodków badawczych w Europie w ramach programu „Cyfrowa Europa”, w razie potrzeby przy wsparciu działań w dziedzinie badań i innowacji w ramach programu „Horyzont Europa”.*

C. UMIEJĘTNOŚCI

Europejskie podejście do sztucznej inteligencji będzie musiało kłaść silny nacisk na umiejętności, aby zapłacić powstający niedobór kompetencji²⁰. Komisja przedstawi wkrótce aktualizację programu na rzecz umiejętności, którego celem jest zapewnienie, by wszyscy mieszkańcy Europy mogli odnieść korzyści z ekologicznej i cyfrowej transformacji gospodarki UE. Inicjatywy mogłyby również obejmować wsparcie sektorowych organów regulacyjnych w zakresie zwiększania ich umiejętności związanych z AI w celu skutecznego i wydajnego wdrażania odpowiednich przepisów. Zaktualizowany plan działania w dziedzinie edukacji cyfrowej przyczyni się do lepszego wykorzystania danych i technologii opartych na sztucznej inteligencji, takich jak uczenie się i analizy predyktywne w celu poprawy systemów kształcenia i szkolenia oraz dostosowania ich do potrzeb ery cyfrowej. Plan zwiększy również powszechną wiedzę na temat sztucznej inteligencji na wszystkich poziomach edukacji, aby przygotować obywateli do podejmowania świadomych decyzji, które będą w coraz większym stopniu opierać się na wykorzystaniu sztucznej inteligencji.

Rozwój umiejętności niezbędnych do pracy z AI i podnoszenie kwalifikacji siły roboczej w celu przygotowania się do transformacji opartej na sztucznej inteligencji będzie priorytetem zrewidowanego skoordynowanego planu w sprawie sztucznej inteligencji, który zostanie opracowany

¹⁹ Przyszły Europejski Fundusz Obronny oraz stała współpraca strukturalna (PESCO) również zapewnią możliwości w zakresie badań i rozwoju w dziedzinie AI. Projekty te powinny być zsynchronizowane z zakrojonymi na szerszą skalę cywilnymi programami UE dotyczącymi sztucznej inteligencji.

²⁰ <https://ec.europa.eu/jrc/en/publication/academic-offer-and-demand-advanced-profiles-eu>

wraz z państwami członkowskimi. Może to obejmować przekształcenie listy kontrolnej zawartej w wytycznych dotyczących etyki w orientacyjny program nauczania dla deweloperów sztucznej inteligencji, który zostanie udostępniony jako zasób instytucjom szkoleniowym. Należy podjąć szczególne starania w celu zwiększenia liczby kobiet przeszkolonych i zatrudnionych w tym obszarze.

Ponadto sztandarowe centrum badań naukowych i innowacji w dziedzinie sztucznej inteligencji w Europie przyciągnęłoby talenty z całego świata z uwagi na możliwości, jakie mogłoby zaoferować. Przyczyniłoby się również do rozwoju i rozpowszechnienia doskonałości w odniesieniu do umiejętności w Europie.

- *Działanie nr 3: Ustanowienie i wspieranie – za pośrednictwem filaru „zaawansowane umiejętności” – sieci wiodących uniwersytetów i instytucji szkolnictwa wyższego w ramach programu „Cyfrowa Europa” w celu przyciągnięcia najlepszych profesorów i naukowców oraz zaoferowania najlepszych na świecie wzorcowych programów magisterskich w zakresie sztucznej inteligencji.*

Obok podnoszenia umiejętności bezpośredni wpływ na pracowników i pracodawców ma architektura i stosowanie systemów sztucznej inteligencji w miejscu pracy. Zaangażowanie partnerów społecznych będzie miało kluczowe znaczenie dla zapewnienia ukierunkowanego na człowieka podejścia do sztucznej inteligencji w miejscu pracy.

D. UKIERUNKOWANIE NA MŚP

Ważne będzie również zapewnienie MŚP dostępu do sztucznej inteligencji i możliwości korzystania z niej. W tym celu należy jeszcze bardziej wzmocnić ośrodki innowacji cyfrowych²¹ i platformę dostępnych na żądanie usług z zakresu sztucznej inteligencji²² oraz wspierać współpracę między MŚP. Program „Cyfrowa Europa” będzie miał tu zasadnicze znaczenie. Chociaż wszystkie ośrodki innowacji cyfrowych powinny zapewniać MŚP wsparcie w zrozumieniu i stosowaniu sztucznej inteligencji, ważne będzie, by co najmniej jeden ośrodek innowacji z każdego państwa członkowskiego posiadał wysoki stopień specjalizacji w dziedzinie sztucznej inteligencji.

MŚP i przedsiębiorstwa typu start-up będą potrzebować dostępu do finansowania w celu dostosowania swoich procedur lub wprowadzania innowacji z wykorzystaniem sztucznej inteligencji. Bazując na planowanym pilotażowym funduszu inwestycyjnym dysponującym środkami w wysokości 100 mln EUR w dziedzinie sztucznej inteligencji i łańcucha bloków, Komisja planuje zwiększyć dostęp do finansowania w zakresie sztucznej inteligencji w ramach programu InvestEU²³. Sztuczna inteligencja została wyraźnie wymieniona wśród kwalifikowalnych obszarów korzystania z gwarancji InvestEU.

- *Działanie nr 4: Komisja będzie współpracować z państwami członkowskimi, aby zapewnić wysoki stopień specjalizacji w zakresie sztucznej inteligencji w co najmniej jednym ośrodku innowacji cyfrowych w każdym państwie członkowskim. Ośrodki innowacji cyfrowych można wspierać w ramach programu „Cyfrowa Europa”.*
- *Komisja i Europejski Fundusz Inwestycyjny uruchomią w pierwszym kwartale 2020 r. program pilotażowy dysponujący budżetem w wysokości 100 mln EUR w celu zapewnienia finansowania kapitałowego na rzecz innowacyjnego rozwoju sztucznej inteligencji.*

²¹ ec.europa.eu/digital-single-market/en/news/digital-innovation-hubs-helping-companies-across-economy-make-most-digital-opportunities.

²² www.Ai4eu.eu.

²³ Europe.eu/investeu

Z zastrzeżeniem ostatecznego porozumienia w sprawie wieloletnich ram finansowych Komisja zamierza znacznie zwiększyć finansowanie od 2021 r. poprzez InvestEU.

E. PARTNERSTWO Z SEKTOREM PRYWATNYM

Należy również zadbać o to, by sektor prywatny był w pełni zaangażowany w określanie programu badań naukowych i innowacji i by zapewniał niezbędny poziom współinwestycji. Wymaga to utworzenia szeroko zakrojonego prywatnego partnerstwa publiczno-prywatnego oraz zaangażowania kadry kierowniczej przedsiębiorstw.

- *Działanie nr 5: W kontekście programu „Horyzont Europa” Komisja ustanowi nowe partnerstwo publiczno-prywatne w dziedzinie AI, danych i robotyki w celu połączenia wysiłków, zapewnienia koordynacji badań naukowych i innowacji w dziedzinie AI, współpracy z innymi partnerstwami publiczno-prywatnymi w ramach programu „Horyzont Europa” oraz współpracy z obiektami badawczymi i ośrodkami innowacji cyfrowych, o których mowa powyżej.*

F. PROMOWANIE STOSOWANIA SZTUCZNEJ INTELIGENCJI PRZEZ SEKTOR PUBLICZNY

Istotne jest, aby administracja publiczna, szpitale, przedsiębiorstwa użyteczności publicznej i przedsiębiorstwa transportowe, organy nadzoru finansowego oraz podmioty działające w innych obszarach interesu publicznego szybko zaczęły stosować produkty i usługi, które w swoim działaniu opierają się na sztucznej inteligencji. Szczególny nacisk zostanie położony na obszary opieki zdrowotnej i transportu, w których technologia jest wystarczająco dojrzała, aby wprowadzić ją na dużą skalę.

- *Działanie nr 6: Komisja zainicjuje otwarte i przejrzyste dialogi sektorowe, traktując priorytetowo opiekę zdrowotną, administracje wiejskie i operatorów usług publicznych, tak aby przedstawić plan działania mający ułatwić rozwój, testy i wprowadzanie AI. Dialogi sektorowe zostaną wykorzystane do przygotowania specjalnego programu służącego wprowadzeniu sztucznej inteligencji, który będzie wspierał zamówienia publiczne na systemy AI i przyczyni się do przekształcenia samych procesów udzielania zamówień publicznych.*

G. ZAPEWNIENIE DOSTĘPU DO DANYCH I INFRASTRUKTURY OBLICZENIOWEJ

Obszary działania przedstawione w niniejszej białej księdze uzupełniają plan przedstawiony równolegle w europejskiej strategii w zakresie danych. Poprawa dostępu do danych i zarządzania nimi jest kwestią o zasadniczym znaczeniu. Bez danych rozwój sztucznej inteligencji i innych zastosowań cyfrowych nie jest możliwy. Ta ogromna ilość nowych danych, które jeszcze nie zostały wygenerowane, stanowi dla Europy szansę na utrzymanie się na czele transformacji w zakresie danych i sztucznej inteligencji. Propagowanie odpowiedzialnych praktyk w zakresie zarządzania danymi i zgodności danych z zasadami FAIR przyczyni się do budowania zaufania i zapewnienia możliwości ponownego wykorzystywania danych²⁴. Równie ważne są inwestycje w kluczowe technologie obliczeniowe i infrastrukturę obliczeniową.

W ramach programu „Cyfrowa Europa” Komisja zaproponowała ponad 4 mld EUR na wspieranie obliczeń wielkiej skali i obliczeń kwantowych, w tym przetwarzania na obrzeżach sieci i sztucznej

²⁴ Dane łatwe do znalezienia, dostępne, interoperacyjne i możliwe do ponownego wykorzystania, jak określono w sprawozdaniu końcowym i planie działania grupy ekspertów Komisji w sprawie danych FAIR, 2018, https://ec.europa.eu/info/sites/info/files/turning_fair_into_reality_1.pdf.

inteligencji, infrastruktury na potrzeby danych i chmury obliczeniowej. W Europejskiej strategii w zakresie danych priorytetom tym poświęcono więcej uwagi.

H. ASPEKTY MIĘDZYNARODOWE

Europa może odgrywać wiodącą rolę w budowaniu sojuszy w oparciu o wspólne wartości oraz w promowaniu etycznego korzystania ze sztucznej inteligencji. Prace UE nad sztuczną inteligencją miały już wpływ na międzynarodowe dyskusje. Opracowując wytyczne etyczne, grupa ekspertów wysokiego szczebla zaangażowała szereg organizacji spoza UE oraz obserwatorów rządowych. Jednocześnie UE była ściśle zaangażowana w rozwój zasad etycznych OECD dotyczących sztucznej inteligencji²⁵. Grupa G-20 poparła te zasady w deklaracji ministrów z czerwca 2019 r. w sprawie handlu i gospodarki cyfrowej.

Jednocześnie UE zauważa, że na innych forach wielostronnych są obecnie prowadzone ważne prace nad sztuczną inteligencją – w tym w Radzie Europy, Organizacji Narodów Zjednoczonych do spraw Oświaty, Nauki i Kultury (UNESCO), Organizacji Współpracy Gospodarczej i Rozwoju (OECD), Światowej Organizacji Handlu oraz Międzynarodowym Związku Telekomunikacyjnym (ITU). W ramach ONZ Unia jest zaangażowana w działania związane ze sprawozdaniem panelu wysokiego szczebla ds. współpracy cyfrowej, w tym z jego zaleceniem dotyczącym sztucznej inteligencji.

UE będzie nadal współpracować z krajami o podobnych poglądach, ale również z globalnymi graczami w dziedzinie sztucznej inteligencji, w oparciu o podejście promujące zasady i wartości UE (np. wspieranie pozytywnej zbieżności przepisów, dostęp do kluczowych zasobów, w tym danych, tworzenie równych szans). Komisja będzie ściśle monitorować strategie polityczne państw trzecich, które ograniczają przepływ danych i zajmie się nadmiernymi ograniczeniami w dwustronnych negocjacjach handlowych oraz poprzez działania w ramach Światowej Organizacji Handlu. Komisja jest przekonana, że współpraca międzynarodowa w dziedzinie AI musi opierać się na podejściu propagującym poszanowanie praw podstawowych, w tym godności ludzkiej, pluralizmu, włączenia, niedyskryminacji i ochrony prywatności i danych osobowych²⁶ i będzie pracować nad szerzeniem tych wartości na całym świecie²⁷. Jasne jest, że odpowiedzialny rozwój i stosowanie sztucznej inteligencji mogą być siłą napędową osiągnięcia celów zrównoważonego rozwoju i realizacji agendy na rzecz zrównoważonego rozwoju 2030.

5. EKOSYSTEM ZAUFANIA: RAMY REGULACYJNE SZTUCZNEJ INTELIGENCJI

Podobnie jak w przypadku każdej nowej technologii wykorzystanie sztucznej inteligencji niesie ze sobą zarówno szanse, jak i zagrożenia. Obywatele obawiają się, że nie będą w stanie bronić swoich praw i bezpieczeństwa, stojąc w obliczu asymetrii informacyjnej związanej z procesem podejmowania decyzji przez algorytmy, zaś przedsiębiorstwa obawiają się niepewności prawa. Chociaż sztuczna inteligencja może przyczynić się do ochrony bezpieczeństwa obywateli i korzystania z przysługujących im praw podstawowych, obywatele obawiają się również, że może ona mieć niezamierzone skutki lub nawet być wykorzystywana w złej wierze. Obawy te wymagają przedstawienia rozwiązań. Co więcej, obok braku inwestycji i umiejętności, głównym czynnikiem powstrzymującym upowszechnianie sztucznej inteligencji jest brak zaufania.

²⁵ <https://www.oecd.org/going-digital/ai/principles/>

²⁶ W ramach Instrumentu Partnerstwa Komisja sfinansuje projekt o wartości 2,5 mln EUR, który ułatwi współpracę z partnerami o podobnych poglądach, w celu promowania unijnych wytycznych etycznych w zakresie sztucznej inteligencji oraz przyjęcia wspólnych zasad i wniosków operacyjnych.

²⁷ Przewodnicząca Von der Leyen, Unia, która mierzy wyżej. Mój program dla Europy, s. 17.

Dlatego też 25 kwietnia 2018 r. Komisja przedstawiła strategię na rzecz sztucznej inteligencji²⁸ uwzględniającą wymiar społeczno-gospodarczy wraz ze zwiększeniem inwestycji w badania naukowe, innowacje i zdolności dotyczące AI w całej Unii. Aby dostosować strategię, z państwami członkowskimi uzgodniono skoordynowany plan²⁹. Komisja powołała również grupę ekspertów wysokiego szczebla, która w kwietniu 2019 r. opublikowała wytyczne w sprawie godnej zaufania sztucznej inteligencji³⁰.

Komisja opublikowała komunikat³¹, w którym z zadowoleniem przyjęła siedem kluczowych wymogów określonych w wytycznych grupy ekspertów wysokiego szczebla:

- przewodnia i nadzorcza rola człowieka,
- techniczna solidność i bezpieczeństwo,
- ochrona prywatności i zarządzanie danymi,
- przejrzystość,
- różnorodność, niedyskryminacja i sprawiedliwość,
- dobrostan społeczny i środowiskowy oraz
- odpowiedzialność.

Ponadto wytyczne zawierają listę kontrolną, która może zostać wykorzystana przez przedsiębiorstwa. W drugiej połowie 2019 r. ponad 350 organizacji przetestowało tę listę i przesłało informacje zwrotne. Grupa wysokiego szczebla jest w trakcie przeglądu wytycznych w kontekście otrzymanych uwag i zakończy te prace do czerwca 2020 r. Z otrzymanych informacji wynika, że chociaż wiele wymogów jest już uwzględnionych w istniejących systemach prawnych lub regulacyjnych, przepisy dotyczące przejrzystości, identyfikowalności i nadzoru przez człowieka nie są uwzględnione w obecnym prawodawstwie w odniesieniu do wielu sektorów gospodarki.

Oprócz tego zbioru niewiążących wytycznych grupy ekspertów wysokiego szczebla oraz zgodnie z wytycznymi politycznymi przewodniczącej jasne europejskie ramy regulacyjne przyczyniłyby się do budowania zaufania do AI wśród konsumentów i przedsiębiorstw, a tym samym przyspieszyłyby upowszechnianie tej technologii. Takie ramy regulacyjne powinny być spójne z innymi działaniami na rzecz wspierania potencjału innowacyjnego i konkurencyjności Europy w tej dziedzinie. Muszą one zapewniać optymalne pod względem społecznym, środowiskowym i gospodarczym wyniki oraz przestrzeganie unijnych przepisów, zasad i wartości. Ma to szczególne znaczenie w obszarach, które bezpośrednio dotyczą praw obywateli, na przykład zastosowań AI w dziedzinach egzekwowania prawa i sądownictwa.

Już dziś deweloperzy i operatorzy sztucznej inteligencji podlegają europejskim przepisom dotyczącym praw podstawowych (np. ochrony danych, prywatności, niedyskryminacji), ochrony konsumentów oraz przepisom w zakresie bezpieczeństwa produktów i odpowiedzialności. Konsumentów oczekują tego samego poziomu bezpieczeństwa i poszanowania ich praw, niezależnie od tego, czy dany produkt lub system opiera się na sztucznej inteligencji. Szczególne cechy sztucznej inteligencji (np. nieprzejrzystość) mogą jednak utrudnić stosowanie i egzekwowanie tych przepisów. Należy więc zbadać, czy obecne przepisy są w stanie uwzględnić ryzyko związane ze sztuczną inteligencją i mogą być skutecznie egzekwowane oraz czy potrzebne jest dostosowanie ustawodawstwa lub nowe przepisy.

²⁸ COM(2018) 237.

²⁹ COM(2018) 795.

³⁰ <https://ec.europa.eu/futurium/en/ai-alliance-consultation/guidelines#Top>

³¹ COM(2019) 168.

Biorąc pod uwagę szybki rozwój sztucznej inteligencji, ramy regulacyjne muszą umożliwiać wprowadzanie dalszych zmian. Wszelkie zmiany powinny ograniczać się do jasno określonych problemów, w przypadku których istnieją wykonalne rozwiązania.

Państwa członkowskie wskazują na obecny brak wspólnych ram europejskich. Niemiecka komisja ds. etyki danych wezwała do opracowania pięciopoziomowego systemu regulacji opartego na ocenie ryzyka – od braku regulacji w odniesieniu do najbardziej nieszkodliwych systemów sztucznej inteligencji po całkowity zakaz dla najbardziej niebezpiecznych systemów. Dania wprowadziła właśnie prototyp pieczęci poświadczającej etyczne wykorzystanie danych. Malta wprowadziła dobrowolny system certyfikacji AI. Jeżeli UE nie zapewni ogólnounijnego podejścia, istnieje realne ryzyko rozdrobnienia rynku wewnętrznego, co podważyłoby cele związane z zaufaniem, pewnością prawa i popularyzacją rozwiązań z zakresu AI na rynku.

Solidne europejskie ramy regulacyjne w zakresie godnej zaufania sztucznej inteligencji będą chronić wszystkich obywateli Unii i pomogą stworzyć bezproblemowy rynek wewnętrzny dla dalszego rozwoju i upowszechniania sztucznej inteligencji, a także wzmocnią bazę przemysłową Europy w dziedzinie sztucznej inteligencji.

A. OKREŚLENIE PROBLEMU

Chociaż sztuczna inteligencja może przynieść wiele korzyści, w tym poprzez zwiększenie bezpieczeństwa produktów i procesów, może ona również wyrządzić szkody. Szkody te mogą być zarówno materialne (dla bezpieczeństwa i zdrowia osób, w tym utrata życia, szkody rzeczowe), jak i niematerialne (utrata prywatności, ograniczenie prawa do wolności słowa, naruszenie godności ludzkiej, dyskryminacja związana np. z dostępem do zatrudnienia) i mogą wiązać się z wieloma różnymi rodzajami ryzyka. Ramy regulacyjne powinny koncentrować się na tym, jak zminimalizować różne rodzaje ryzyka związane z potencjalnymi szkodami, zwłaszcza najpoważniejszymi.

Główne zagrożenia związane z AI dotyczą stosowania przepisów mających na celu ochronę praw podstawowych (w tym przepisów dotyczących ochrony danych osobowych i prywatności oraz niedyskryminacji), a także kwestii związanych z bezpieczeństwem³² i odpowiedzialnością.

Zagrożenia dla praw podstawowych, w tym w odniesieniu do ochrony danych osobowych i prywatności oraz niedyskryminacji

Wykorzystywanie sztucznej inteligencji może mieć wpływ na wartości, na których zbudowana jest Unia i prowadzić do naruszenia praw podstawowych³³, w tym prawa do wolności wypowiedzi, wolności zgromadzeń, godności ludzkiej, braku dyskryminacji ze względu na płeć, rasę, pochodzenie etniczne, religię lub wierzenia, niepełnosprawność, wiek lub orientację seksualną, ochrony danych osobowych i życia prywatnego³⁴ lub prawa do dochodzenia odszkodowania na drodze sądowej i rzetelnego procesu, jak również ochrony konsumentów. Zagrożenia te mogą być wynikiem

³² Obejmuje to kwestie cyberbezpieczeństwa, kwestie związane z zastosowaniami AI w infrastrukturze krytycznej lub wykorzystywanie sztucznej inteligencji w złej wierze.

³³ Z badań Rady Europy wynika, że korzystanie ze sztucznej inteligencji może mieć wpływ na wiele praw podstawowych: <https://rm.coe.int/algorithms-and-human-rights-en-rev/16807956b5>

³⁴ Ogólne rozporządzenie o ochronie danych i dyrektywa o prywatności i łączności elektronicznej (nowe rozporządzenie w sprawie prywatności i łączności elektronicznej w trakcie negocjacji) regulują tego rodzaju zagrożenia, ale może zaistnieć potrzeba zbadania, czy systemy sztucznej inteligencji stwarzają dodatkowe ryzyko. Komisja będzie monitorować i oceniać stosowanie RODO w sposób ciągły.

wadliwego opracowania systemów sztucznej inteligencji (w tym w odniesieniu do nadzoru przez człowieka) lub wykorzystywania danych bez korygowania możliwych uprzedzeń (np. system jest szkolony z wykorzystaniem wyłącznie lub głównie danych pochodzących od mężczyzn, co może prowadzić do nieoptymalnych wyników w odniesieniu do kobiet).

Sztuczna inteligencja może pełnić wiele funkcji, które wcześniej mogły być wykonywane wyłącznie przez człowieka. W związku z tym obywatele i osoby prawne będą w coraz większym stopniu podlegać działaniom i decyzjom podejmowanym przez systemy sztucznej inteligencji lub z ich pomocą, które czasami mogą być trudne do zrozumienia i od których może być się trudno odwołać w razie potrzeby. Ponadto sztuczna inteligencja zwiększa możliwość śledzenia i analizowania codziennych poczynąń ludzi. Na przykład istnieje potencjalne ryzyko, że sztuczna inteligencja może być wykorzystywana – z naruszeniem unijnych przepisów dotyczących ochrony danych i innych – przez organy państwowe lub inne podmioty prowadzące masowy nadzór oraz przez pracodawców obserwujących zachowanie pracowników. Poprzez dokonywanie analizy dużych ilości danych i identyfikację powiązań między nimi sztuczna inteligencja może być również wykorzystywana do znajdowania źródeł pochodzenia danych i deanonimizacji danych dotyczących osób, co stwarza nowe zagrożenie dla ochrony danych osobowych nawet w odniesieniu do zbiorów danych, które same w sobie nie obejmują danych osobowych. Sztuczna inteligencja jest również wykorzystywana przez pośredników internetowych do priorytetowego traktowania informacji dla swoich użytkowników i moderacji treści. Przetwarzane dane, projekt aplikacji i możliwość interwencji człowieka mogą mieć wpływ na prawo do wolności słowa, ochrony danych osobowych, prywatności i swobód politycznych.

Niektóre algorytmy sztucznej inteligencji, wykorzystywane do przewidywania recydywy, mogą wykazywać dyskryminację ze względu na płeć i rasę, wskazując różne prawdopodobieństwo recydywy dla kobiet i mężczyzn lub dla obywateli i cudzoziemców. Źródło: Tolan S., Miron M., Gomez E. and Castillo C. „*Why Machine Learning May Lead to Unfairness: Evidence from Risk Assessment for Juvenile Justice in Catalonia*” (Dlaczego uczenie maszynowe może prowadzić do braku sprawiedliwości: dowody pochodzące z oceny ryzyka w odniesieniu do wymiaru sprawiedliwości dla nieletnich w Katalonii), Best Paper Award, International Conference on AI and Law, 2019.

Niektóre programy sztucznej inteligencji służące do analizy twarzy wykazują dyskryminację ze względu na płeć lub rasę, charakteryzując się niskim poziomem błędów przy określaniu płci mężczyzn o jaśniejszym kolorze skóry i wysokim poziomem błędów przy określaniu płci kobiet o ciemniejszym kolorze skóry. Źródło: Joy Buolamwini, Timnit Gebru; *Proceedings of the 1st Conference on Fairness, Accountability and Transparency (Ustalenia 1. konferencji na temat uczciwości, odpowiedzialności i przejrzystości)*, PMLR 81:77-91, 2018.

Uprzedzenia i dyskryminacja stanowią nieodłączne ryzyko wszelkiej działalności społecznej lub gospodarczej. Proces podejmowania decyzji przez człowieka nie jest odporny na błędy i uprzedzenia. Jednak takie samo uprzedzenie, które występuje w przypadku sztucznej inteligencji, może mieć znacznie większe konsekwencje i prowadzić do dyskryminacji znacznej liczby osób ze względu na

brak mechanizmów kontroli społecznej zbliżonych do tych kierujących zachowaniami ludzi³⁵. Może się to zdarzyć również wtedy, gdy system „uczy się” w trakcie działania. W takich przypadkach, gdy nie można było przewidzieć danego wyniku lub mu zapobiec na etapie projektowania, ryzyko nie wynika z wady pierwotnego projektu systemu, lecz z praktycznych skutków korelacji lub wzorców, które system identyfikuje w dużym zbiorze danych.

Szczególne cechy technologii sztucznej inteligencji, w tym nieprzejrzystość („efekt czarnej skrzynki”), złożoność, nieprzewidywalność i częściowo samodzielne działanie mogą utrudniać weryfikację zgodności oraz egzekwowanie istniejących przepisów UE służących ochronie praw podstawowych. Organy egzekwowania prawa i zainteresowane osoby mogą nie mieć możliwości sprawdzenia, w jaki sposób dana decyzja, w którą zaangażowana była AI, została podjęta, a zatem czy obowiązujące przepisy były przestrzegane. Osoby fizyczne i prawne mogą mieć trudności z faktycznym dostępem do wymiaru sprawiedliwości w sytuacjach, w których takie decyzje mogą mieć na nie negatywny wpływ.

Zagrożenia dla bezpieczeństwa i skutecznego funkcjonowania systemu odpowiedzialności

Technologie sztucznej inteligencji wbudowane w produkty i usługi mogą stwarzać dla użytkowników nowe zagrożenia związane z bezpieczeństwem. Przykładowo, w wyniku błędu w technologii rozpoznawania obiektu, samochód autonomiczny może błędnie zidentyfikować obiekt na drodze i spowodować wypadek, w którym dojdzie do obrażeń ciała i szkód materialnych. Podobnie jak w przypadku zagrożeń dla praw podstawowych, zagrożenia te mogą być spowodowane wadami w projekcie technologii sztucznej inteligencji, być związane z problemami z dostępnością i jakością danych lub z innymi problemami wynikającymi z uczenia się maszyn. Chociaż niektóre z tych zagrożeń nie są ograniczone do produktów i usług, które opierają się na sztucznej inteligencji, wykorzystanie sztucznej inteligencji może spowodować ich nasilenie.

Brak jasnych przepisów mających zaradzić tym zagrożeniom może – obok zagrożeń dla zainteresowanych osób – prowadzić do braku pewności prawa dla przedsiębiorstw, które wprowadzają do obrotu w UE swoje produkty wykorzystujące sztuczną inteligencję. Organy nadzoru rynku i egzekwowania prawa mogą znaleźć się w sytuacji, w której nie wiedzą, czy mogą interweniować, ponieważ nie są uprawnione do działania lub nie dysponują odpowiednimi zdolnościami technicznymi, aby skontrolować systemy³⁶. Niepewność prawa może zatem obniżyć ogólny poziom bezpieczeństwa i osłabić konkurencyjność europejskich przedsiębiorstw.

W przypadku urzeczywistnienia się zagrożeń dla bezpieczeństwa, brak jasnych wymogów oraz wymienione powyżej cechy technologii AI sprawiają, że trudno będzie prześledzić potencjalnie problematyczne decyzje podejmowane z wykorzystaniem systemów AI. To z kolei może utrudniać

³⁵ Komitet Doradczy ds. Równości Szans dla Kobiet i Mężczyzn przygotowuje obecnie „opinię w sprawie sztucznej inteligencji” zawierającą między innymi analizę wpływu sztucznej inteligencji na równość płci. Opinia ma zostać przyjęta przez komitet na początku 2020 r. Europejska strategia na rzecz równouprawnienia płci na lata 2020–2024 również odnosi się do związku między sztuczną inteligencją a równouprawnieniem płci. Na początku 2020 r. europejska sieć krajowych organów ds. równości (Equinet) opublikuje sprawozdanie (opracowane przez Robina Allena i Dee Masters) zatytułowane „Regulating AI: the new role for Equality Bodies – Meeting the new challenges to equality and non-discrimination from increased digitalisation and the use of AI” (Regulacja w zakresie AI: nowa rola organów ds. równości – Sprostanie nowym wyzwaniom w zakresie równości i niedyskryminacji wynikającym z rosnącej cyfryzacji i stosowania sztucznej inteligencji).

³⁶ Przykładem może być inteligentny zegarek dla dzieci. Produkt ten nie wyrządziłby bezpośredniej szkody dziecku noszącemu ten zegarek, ale ponieważ nie posiada on minimalnego poziomu zabezpieczeń, może być łatwo wykorzystany jako narzędzie dostępu do dziecka. Organy nadzoru rynku mogą mieć trudności z interwencją w przypadkach, gdy ryzyko nie jest powiązane z produktem jako takim.

poszkodowanym uzyskanie odszkodowania na mocy obowiązującego prawodawstwa unijnego i krajowego w dziedzinie odpowiedzialności³⁷.

Na mocy dyrektywy w sprawie odpowiedzialności za produkty producent odpowiada za szkodę wyrządzoną przez wadliwy produkt. Jednak w przypadku systemów opartych na AI, takich jak samochody autonomiczne, może być trudno udowodnić, że produkt jest wadliwy, udowodnić wyrządzoną szkodę lub też związek przyczynowy między nimi. Ponadto nie do końca wiadomo, jak i w jakim stopniu dyrektywę w sprawie odpowiedzialności za produkty stosuje się w przypadku pewnych rodzajów wad, na przykład wynikających z niedociągnięć związanych z cyberbezpieczeństwem produktu.

Wyzwania związane z trudnościami w śledzeniu potencjalnie problematycznych decyzji podejmowanych przez systemy sztucznej inteligencji, o których mowa powyżej w odniesieniu do praw podstawowych, stosują się w równym stopniu do kwestii bezpieczeństwa i odpowiedzialności. Osoby poszkodowane mogą nie mieć skutecznego dostępu do dowodów, które są niezbędne do przeprowadzenia postępowania w sądzie i mogą mieć mniejsze możliwości uzyskania odszkodowania w porównaniu z sytuacjami, w których szkoda jest spowodowana przez tradycyjne technologie. Takie ryzyko będzie wzrastać w miarę jak wykorzystanie sztucznej inteligencji stawać się będzie coraz powszechniejsze.

B. EWENTUALNE DOSTOSOWANIE DO ISTNIEJĄCYCH RAM PRAWNYCH UE W DZIEDZINIE SZTUCZNEJ INTELIGENCJI

Obszerny zbiór istniejących unijnych przepisów w zakresie bezpieczeństwa i odpowiedzialności w odniesieniu do produktów³⁸, w tym przepisów sektorowych, uzupełniony przepisami krajowymi, ma znaczenie i może być stosowany do szeregu nowych zastosowań w dziedzinie sztucznej inteligencji.

Jeżeli chodzi o ochronę praw podstawowych i praw konsumentów, unijne ramy prawne obejmują takie przepisy jak dyrektywa w sprawie równości rasowej³⁹, dyrektywa w sprawie równego traktowania w zakresie zatrudnienia i pracy⁴⁰, dyrektywy w sprawie równego traktowania kobiet i mężczyzn w odniesieniu do zatrudnienia i dostępu do towarów i usług⁴¹, szereg przepisów dotyczących ochrony konsumentów⁴², a także przepisy dotyczące ochrony danych osobowych i prywatności, w szczególności ogólne rozporządzenie o ochronie danych i inne przepisy sektorowe dotyczące ochrony danych osobowych, takie jak dyrektywa o ochronie danych w sprawach karnych⁴³. Od 2025 r. obowiązywać będą przepisy dotyczące wymogów w zakresie dostępności towarów i usług określone w europejskim akcie w sprawie dostępności⁴⁴. Ponadto przy wdrażaniu innych przepisów unijnych,

³⁷ Wpływ sztucznej inteligencji, internetu rzeczy i innych technologii cyfrowych na przepisy dotyczące bezpieczeństwa i odpowiedzialności został przeanalizowany w sprawozdaniu Komisji towarzyszącym niniejszej białej księdze.

³⁸ Unijne ramy prawne dotyczące bezpieczeństwa produktów obejmują dyrektywę w sprawie ogólnego bezpieczeństwa produktów (2001/95/WE) oraz szereg przepisów sektorowych obejmujących różne kategorie produktów, od maszyn, samolotów i samochodów po zabawki i wyroby medyczne, których celem jest zapewnienie wysokiego poziomu zdrowia i bezpieczeństwa. Uzupełnieniem prawa dotyczącego odpowiedzialności za produkt są różne systemy odpowiedzialności cywilnej za szkody spowodowane przez produkty lub usługi.

³⁹ Dyrektywa 2000/43/WE.

⁴⁰ Dyrektywa 2000/78/WE.

⁴¹ Dyrektywa 2004/113/WE; dyrektywa 2006/54/WE.

⁴² Jak np. dyrektywa w sprawie nieuczciwych praktyk handlowych (2005/29/WE) i dyrektywa w sprawie praw konsumentów (2011/83/WE).

⁴³ Dyrektywa Parlamentu Europejskiego i Rady (UE) 2016/680 z dnia 27 kwietnia 2016 r. w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych przez właściwe organy do celów zapobiegania przestępczości, prowadzenia postępowań przygotowawczych, wykrywania i ścigania czynów zabronionych i wykonywania kar, oraz w sprawie swobodnego przepływu takich danych.

⁴⁴ Dyrektywa (UE) 2019/882 w sprawie wymogów dostępności produktów i usług.

w tym w dziedzinie usług finansowych, migracji lub odpowiedzialności pośredników internetowych, należy przestrzegać praw podstawowych.

Chociaż prawodawstwo UE zasadniczo nadal ma pełne zastosowanie niezależnie od stopnia wykorzystania AI, ważne jest, aby ocenić, czy można je skutecznie egzekwować w celu wyeliminowania ryzyka, jakie stwarzają systemy sztucznej inteligencji, lub czy konieczne jest dostosowanie konkretnych instrumentów prawnych.

Na przykład podmioty gospodarcze pozostają w pełni odpowiedzialne za zgodność sztucznej inteligencji z obowiązującymi przepisami, które chronią konsumentów, zaś wykorzystanie algorytmów do profilowania konsumentów z naruszeniem odpowiednich przepisów nie jest dozwolone, a naruszenia będą odpowiednio karane.

Komisja jest zdania, że można ulepszyć ramy legislacyjne, aby uwzględnić następujące zagrożenia i sytuacje:

- *Skuteczne stosowanie i egzekwowanie obowiązujących przepisów unijnych i krajowych:* najważniejsze cechy sztucznej inteligencji stanowią wyzwanie dla zapewnienia właściwego stosowania i egzekwowania przepisów unijnych i krajowych. Brak przejrzystości AI sprawia, że trudno jest określić i udowodnić ewentualne naruszenia prawa, w tym przepisów, które chronią prawa podstawowe, przypisują odpowiedzialność i określają warunki konieczne do dochodzenia odszkodowania. W związku z tym, aby zapewnić skuteczne stosowanie i egzekwowanie przepisów, konieczne może być dostosowanie lub doprecyzowanie obowiązujących przepisów w niektórych obszarach, na przykład w odniesieniu do odpowiedzialności, jak szczegółowo opisano w sprawozdaniu towarzyszącym niniejszej białej księdze.
- *Ograniczenia zakresu obowiązującego prawodawstwa UE:* prawodawstwo UE w zakresie bezpieczeństwa produktów koncentruje się przede wszystkim na wprowadzaniu produktów do obrotu. Chociaż w prawodawstwie UE dotyczącym bezpieczeństwa produktów przewidziano, że oprogramowanie stanowiące część końcowego produktu musi być zgodne z odpowiednimi przepisami dotyczącymi bezpieczeństwa produktów, to kwestią otwartą pozostaje, czy samodzielne oprogramowanie jest objęte prawodawstwem UE w zakresie bezpieczeństwa produktów poza pewnymi sektorami o wyraźnie określonych zasadach⁴⁵. Ogólne przepisy unijne dotyczące bezpieczeństwa obowiązujące obecnie mają zastosowanie do produktów, a nie do usług, a zatem zasadniczo nie mają zastosowania do usług opartych na sztucznej inteligencji (np. usług zdrowotnych, finansowych, transportowych).
- *Zmieniająca się funkcjonalność systemów sztucznej inteligencji:* włączenie oprogramowania, w tym sztucznej inteligencji, do produktów może zmienić funkcjonowanie takich produktów i systemów w trakcie ich cyklu życia. Dotyczy to w szczególności systemów, które wymagają częstych aktualizacji oprogramowania lub opierają się na uczeniu maszynowym. Cechy te mogą wiązać się z nowymi zagrożeniami, które nie występowały w momencie wprowadzenia systemu na rynek. Zagrożenia te nie są odpowiednio uwzględnione w obowiązującym prawodawstwie, które koncentruje się głównie na ryzyku związanym z bezpieczeństwem produktów w momencie wprowadzania ich do obrotu.

⁴⁵ Na przykład oprogramowanie przeznaczone przez producenta do stosowania w celach medycznych uznaje się za wyrób medyczny zgodnie z rozporządzeniem w sprawie wyrobów medycznych (UE) 2017/745.

- *Niepewność dotycząca podziału obowiązków między różnymi podmiotami gospodarczymi w łańcuchu dostaw:* ogólnie rzecz biorąc, prawodawstwo UE w zakresie bezpieczeństwa produktów nakłada odpowiedzialność na producenta produktu wprowadzonego do obrotu, w tym za wszystkie jego elementy, np. systemy sztucznej inteligencji. Przepisy mogą stać się jednak niejasne, na przykład w przypadku dodania sztucznej inteligencji przez podmiot, który nie jest producentem, po wprowadzeniu produktu na rynek. Ponadto przepisy UE dotyczące odpowiedzialności za produkty przewidują odpowiedzialność producentów i pozostawiają kwestię uregulowania odpowiedzialności innych podmiotów w łańcuchu dostaw krajowym przepisom dotyczącym odpowiedzialności.
- *Zmiany pojęcia bezpieczeństwa:* wykorzystywanie sztucznej inteligencji w produktach i usługach może stwarzać zagrożenia, które obecnie nie są przedmiotem przepisów UE. Zagrożenia te mogą dotyczyć cyberbezpieczeństwa, bezpieczeństwa osobistego (na przykład w związku z nowymi zastosowaniami sztucznej inteligencji, takimi jak w urządzeniach gospodarstwa domowego), utraty łączności itd. Zagrożenia te mogą występować w momencie wprowadzania produktów do obrotu lub powstawać w wyniku aktualizacji oprogramowania lub uczenia się maszyn w trakcie stosowania produktu. UE powinna w pełni wykorzystać dostępne jej narzędzia, aby zwiększyć bazę dowodową na temat potencjalnych zagrożeń związanych z zastosowaniami AI, w tym wykorzystać doświadczenia Agencji UE ds. Cyberbezpieczeństwa (ENISA) do oceny krajobrazu zagrożeń związanych ze sztuczną inteligencją.

Jak już wspomniano, kilka państw członkowskich analizuje już możliwości związane z prawodawstwem krajowym w celu sprostania wyzwaniom stwarzanym przez sztuczną inteligencję. Zwiększa to ryzyko rozdrobnienia jednolitego rynku. Rozbieżne przepisy krajowe mogą stwarzać przeszkody dla przedsiębiorstw, które chcą sprzedawać i eksploatować systemy sztucznej inteligencji na jednolitym rynku. Zapewnienie wspólnego podejścia na szczeblu UE przyniosłoby europejskim przedsiębiorstwom korzyści wynikające ze sprawnego dostępu do jednolitego rynku i zwiększyłoby ich konkurencyjność na rynkach światowych.

Sprawozdanie na temat wpływu sztucznej inteligencji, internetu rzeczy i robotyki na bezpieczeństwo i odpowiedzialność

Sprawozdanie, które towarzyszy niniejszej białej księdze, zawiera analizę odpowiednich ram prawnych. Określono w nim niejasności związane ze stosowaniem tych ram w odniesieniu do konkretnych zagrożeń związanych ze sztuczną inteligencją i innymi technologiami cyfrowymi.

Stwierdzono, że obecne przepisy dotyczące bezpieczeństwa produktów już teraz wspierają rozszerzoną koncepcję bezpieczeństwa chroniącą przed wszelkiego rodzaju ryzykiem stwarzanym przez produkt w zależności od jego zastosowania. W celu zapewnienia większej pewności prawa można jednak wprowadzić przepisy wyraźnie obejmujące nowe zagrożenia związane z powstającymi technologiami cyfrowymi.

- Autonomiczne zachowanie niektórych systemów sztucznej inteligencji w trakcie ich cyklu życia może pociągać za sobą istotne zmiany w produkcji mające wpływ na bezpieczeństwo, co może wymagać nowej oceny ryzyka. Ponadto – jako zabezpieczenie – może być niezbędny nadzór ze strony człowieka od etapu projektowania produktu i przez cały cykl życia produktów i systemów wykorzystujących AI.
- W stosownych przypadkach można również rozważyć nałożenie na producentów wyraźnych obowiązków w odniesieniu do ryzyka dla bezpieczeństwa psychicznego użytkowników (np. współpraca z robotami humanoidalnymi).
- Unijne przepisy dotyczące bezpieczeństwa produktów mogłyby zawierać szczegółowe wymogi w zakresie zapobiegania zagrożeniu dla bezpieczeństwa, jakie wiąże się z błędnymi danymi na etapie koncepcji, oraz mechanizmy zapewniające utrzymanie jakości danych w trakcie całego okresu stosowania produktów i systemów opartych na AI.
- Brak przejrzystości systemów opartych na algorytmach można rozwiązać za pomocą wymogów w zakresie przejrzystości.
- Jeżeli samodzielne oprogramowanie – wprowadzane do obrotu jako oddzielny produkt lub wgrywane do produktu po wprowadzeniu tego produktu do obrotu – ma wpływ na bezpieczeństwo, konieczne może być dostosowanie lub doprecyzowanie istniejących przepisów.
- Biorąc pod uwagę coraz większą złożoność łańcuchów dostaw w odniesieniu do nowych technologii, przepisy wymagające konkretnie współpracy między podmiotami gospodarczymi w łańcuchu dostaw i użytkownikami mogłyby zapewnić pewność prawa.

Pewne cechy nowych technologii cyfrowych, takich jak AI, IoT czy robotyka, mogą podważać pewne aspekty tych ram odpowiedzialności i ograniczyć ich skuteczność. Niektóre z tych cech mogą utrudnić prześledzenie ścieżki prowadzącej od ludzkiego zachowania do powstania szkody, której ustalenie byłoby niezbędne do dochodzenia roszczenia na zasadzie winy zgodnie z większością przepisów krajowych. Mogłoby to znacznie zwiększyć koszty dla poszkodowanych i oznaczać, że roszczenia odszkodowawcze wobec innych podmiotów niż producenci mogą być trudne do zrealizowania lub udowodnienia.

- Osobom, które poniosły szkodę w wyniku działania systemów sztucznej inteligencji, musi przysługiwać taki sam poziom ochrony co osobom, które doznały szkody spowodowanej innymi technologiami, przy czym należy umożliwić dalszy rozwój innowacji technologicznych.
- Należy starannie ocenić wszystkie możliwości osiągnięcia tego celu, w tym ewentualne zmiany w dyrektywie w sprawie odpowiedzialności za produkty oraz możliwość dalszej ukierunkowanej harmonizacji przepisów krajowych dotyczących odpowiedzialności. Przykładowo Komisja pragnie zasięgnąć opinii co do tego, czy i w jakim stopniu konieczne może być złagodzenie skutków złożoności przez dostosowanie ciężaru dowodu wymaganego na mocy krajowych przepisów dotyczących odpowiedzialności za szkody spowodowane eksploatacją zastosowań opartych na AI.

Na podstawie powyższej dyskusji Komisja stwierdza, że oprócz ewentualnych zmian w obowiązującym prawodawstwie konieczne może być wprowadzenie nowych przepisów dotyczących konkretnie sztucznej inteligencji w celu dostosowania unijnych ram prawnych do aktualnych i oczekiwanych zmian technologicznych i handlowych.

C. ZAKRES PRZYSZŁYCH RAM PRAWNYCH UE

Ważną kwestią dotyczącą przyszłych ram regulacyjnych dotyczących AI jest określenie zakresu ich stosowania. Roboczo zakłada się, że ramy regulacyjne miałyby zastosowanie do produktów i usług opartych na sztucznej inteligencji. W związku z tym sztuczna inteligencja powinna zostać jasno zdefiniowana do celów niniejszej białej księgi oraz wszelkich przyszłych inicjatyw politycznych.

W swoim komunikacie w sprawie sztucznej inteligencji dla Europy Komisja przedstawiła pierwszą definicję sztucznej inteligencji⁴⁶. Definicja ta została doprecyzowana przez grupę ekspertów wysokiego szczebla⁴⁷.

Definicja sztucznej inteligencji w każdym nowym instrumencie prawnym będzie musiała być wystarczająco elastyczna, aby uwzględnić postęp techniczny, a jednocześnie wystarczająco dokładna, aby zapewnić niezbędną pewność prawa.

Do celów niniejszej białej księgi, a także wszelkich możliwych przyszłych dyskusji na temat inicjatyw politycznych, istotne wydaje się wyjaśnienie głównych elementów składających się na sztuczną inteligencję, którymi są „dane” i „algorytmy”. Sztuczna inteligencja może być wbudowana w sprzęt komputerowy. W przypadku technik uczenia maszynowego, które stanowią podzbiór AI, algorytmy zostały wyszkolone do wyprowadzania pewnych wzorców na

Na przykład w przypadku jazdy autonomicznej algorytm wykorzystuje w czasie rzeczywistym dane pochodzące z samochodu (prędkość, zużycie paliwa, amortyzatory itd.) oraz z czujników skanujących otoczenie samochodu (drogę, znaki drogowe, inne pojazdy, pieszych itd.), aby ustalić, jaki kierunek, przyspieszenie i prędkość wybrać w celu dotarcia do wybranego miejsca. Algorytm uczy się na podstawie obserwowanych danych i dostosowuje się do sytuacji na drodze i warunków zewnętrznych, w tym zachowania innych kierowców, aby wybrać najbardziej komfortowy i najbezpieczniejszy sposób jazdy.

podstawie zbioru danych, aby określić działania niezbędne do osiągnięcia danego celu. Algorytmy mogą nadal uczyć się w trakcie działania. Produkty oparte na AI mogą wprawdzie działać autonomicznie przez postrzeganie swojego środowiska i bez konieczności stosowania z góry określonego zestawu instrukcji, ich zachowanie jest jednak w dużej mierze zdefiniowane i ograniczone przez deweloperów. To, do jakich celów system powinien zostać zoptymalizowany, określa i programuje człowiek.

UE dysponuje rygorystycznymi ramami prawnymi zapewniającymi m.in. ochronę konsumentów, służącymi rozwiązaniu problemu nieuczciwych praktyk handlowych i ochronie danych osobowych i prywatności. Ponadto unijny dorobek prawny zawiera szczegółowe przepisy dotyczące niektórych sektorów (np. opieki zdrowotnej, transportu). Te obowiązujące przepisy prawa UE będą miały nadal zastosowanie do sztucznej inteligencji, chociaż konieczne mogą być ich aktualizacje, aby

⁴⁶ COM(2018) 237 final, s. 1: „Termin sztuczna inteligencja odnosi się do systemów, które wykazują inteligentne zachowanie dzięki analizie otoczenia i podejmowaniu działań – do pewnego stopnia autonomicznie – w celu osiągnięcia konkretnych celów. Systemy SI mogą być oparte na oprogramowaniu, działając w świecie wirtualnym (np. asystenci głosowi, oprogramowanie do analizy obrazu, wyszukiwarki, systemy rozpoznawania mowy i twarzy), lub mogą być wbudowane w urządzenia (np. zaawansowane roboty, samochody autonomiczne, drony lub aplikacje internetu rzeczy)”.

⁴⁷ Grupa ekspertów wysokiego szczebla, definicja AI, s. 8: „Systemy sztucznej inteligencji (AI) to oprogramowanie (i ewentualnie również sprzęt komputerowy) zaprojektowane przez człowieka, które – aby osiągnąć złożony cel – działa w wymiarze fizycznym lub cyfrowym, postrzegając swoje środowisko poprzez pozyskiwanie danych, interpretując zgromadzone dane (ustrukturyzowane lub nie), wyciągając wnioski na podstawie tych danych lub przetwarzając informacje, których źródłem są te dane oraz podejmując decyzje w sprawie najlepszych działań, jakie należy podjąć, aby zrealizować ten cel. Systemy sztucznej inteligencji mogą wykorzystywać zasady symboliczne albo uczyć się na podstawie modelu numerycznego i mogą również dostosować swoje zachowanie poprzez analizę wpływu ich wcześniejszych działań na środowisko.”

odzwierciedlić transformację cyfrową i wykorzystanie AI (zob. sekcja B). W związku z tym kwestie, do których odnoszą się już istniejące przepisy horyzontalne lub sektorowe (np. dotyczące wyrobów medycznych⁴⁸ czy systemów transportowych) będą nadal regulowane przez te przepisy.

Nowe ramy regulacyjne w zakresie sztucznej inteligencji powinny co do zasady być wystarczająco skuteczne, aby móc osiągnąć swoje cele, jednak nie powinny być nadmiernie nakazowe, co mogłoby prowadzić do nieproporcjonalnych obciążeń, zwłaszcza dla MŚP. Zdaniem Komisji osiągnięcie tej równowagi będzie możliwe dzięki przyjęciu podejścia opartego na analizie ryzyka.

Podejście oparte na analizie ryzyka ma istotne znaczenie dla zapewnienia proporcjonalności interwencji regulacyjnej. Wymaga to jednak jasnych kryteriów pozwalających na rozróżnienie między zastosowaniami sztucznej inteligencji, w szczególności w odniesieniu do kwestii, czy charakteryzują się one „wysokim ryzykiem”⁴⁹. Określenie, co stanowi zastosowanie wysokiego ryzyka, powinno być jasne i łatwe do zrozumienia oraz możliwe do zastosowania przez wszystkie zainteresowane strony. Jednak nawet jeżeli zastosowanie AI nie kwalifikuje się jako zastosowanie wysokiego ryzyka, jest ono w pełni objęte już obowiązującymi przepisami unijnymi.

Komisja jest zdania, że dane zastosowanie AI należy zasadniczo uznawać za charakteryzujące się wysokim ryzykiem w zależności od potencjalnych konsekwencji danego zastosowania, na podstawie analizy tego, czy zarówno sektor, jak i planowane zastosowanie wiążą się ze znaczącym ryzykiem, w szczególności z punktu widzenia ochrony bezpieczeństwa, praw konsumentów i praw podstawowych. W szczególności należy uznać, że dane zastosowanie sztucznej inteligencji jest zastosowaniem wysokiego ryzyka, jeżeli spełnia łącznie następujące dwa kryteria:

- po pierwsze zastosowanie sztucznej inteligencji dotyczy sektora, w którym, biorąc pod uwagę cechy charakterystyczne zazwyczaj podejmowanych działań, można oczekiwać wystąpienia znaczącego ryzyka. To pierwsze kryterium gwarantuje, że interwencja regulacyjna jest ukierunkowana na obszary, w których – ogólnie rzecz biorąc – ryzyko jest najbardziej prawdopodobne. Przedmiotowe sektory powinny zostać wymienione w szczegółowym i wyczerpującym wykazie w nowych ramach regulacyjnych. Są to na przykład: opieka zdrowotna, transport, energetyka i części sektora publicznego⁵⁰. Wykaz powinien być okresowo poddawany przeglądowi i w razie konieczności zmieniany w zależności od istotnych zmian praktyki;
- po drugie rozwiązanie z zakresu sztucznej inteligencji w danym sektorze jest ponadto stosowane w taki sposób, że istnieje prawdopodobieństwo wystąpienia znaczącego ryzyka. To drugie kryterium odzwierciedla fakt, że nie każde użycie sztucznej inteligencji w wybranych sektorach musi wiązać się ze znaczącym ryzykiem. Na przykład, chociaż sektor opieki zdrowotnej na ogół może być potencjalnie źródłem wysokiego ryzyka, błąd w systemie umawiania wizyt w szpitalu zazwyczaj nie stwarza ryzyka o znaczeniu uzasadniającym interwencję legislacyjną. Ocena stopnia ryzyka danego zastosowania może opierać się na wpływie na zainteresowane strony. Przykładowo chodzi tu o zastosowania AI, które mogą wywierać skutki prawne lub podobnie istotne skutki w odniesieniu do praw jednostki lub

⁴⁸ Na przykład zachodzą różne względy bezpieczeństwa i powstają skutki prawne dotyczące systemów sztucznej inteligencji, które zapewniają lekarzom specjalistyczne informacje medyczne, systemów dostarczających informacji medycznych bezpośrednio pacjentom i systemów wykonujących zadania medyczne bezpośrednio na pacjencie. Komisja bada obecnie te wyzwania związane z bezpieczeństwem i odpowiedzialnością, które są specyficzne dla opieki zdrowotnej.

⁴⁹ W prawodawstwie UE „ryzyko” może zostać sklasyfikowane w inny sposób niż w niniejszym dokumencie, w zależności od obszaru, na przykład w odniesieniu do bezpieczeństwa produktów.

⁵⁰ Sektor publiczny mógłby obejmować takie obszary jak polityka azylowa, migracja, kontrole graniczne, sądownictwo, zabezpieczenie społeczne i służby zatrudnienia.

przedsiębiorstwa; które stwarzają ryzyko uszkodzenia ciała, śmierci lub istotnej szkody materialnej lub niematerialnej; lub których skutków osoby fizyczne lub prawne nie mogą w rozsądny sposób uniknąć.

Zastosowanie tych dwóch łącznych kryteriów zapewniłoby ukierunkowanie zakresu ram regulacyjnych i zagwarantowałoby pewność prawa. Obowiązkowe wymogi zawarte w nowych ramach regulacyjnych dotyczących sztucznej inteligencji (zob. sekcja D poniżej) miałyby zasadniczo zastosowanie wyłącznie do tych zastosowań, które uznano za charakteryzujące się wysokim ryzykiem zgodnie z tymi dwoma łącznymi kryteriami.

Niezależnie od powyższego mogą również wystąpić wyjątkowe przypadki, w których ze względu na ryzyko związane z wykorzystaniem AI do określonych celów należy uznać je za zastosowania wysokiego ryzyka – to znaczy niezależnie od sektora i w przypadku gdy nadal zastosowanie mają poniższe wymogi⁵¹. Przykładowo:

- wykorzystanie AI w procesie rekrutacji oraz w sytuacjach mających wpływ na prawa pracowników zawsze byłoby uznawane za charakteryzujące się wysokim ryzykiem, a zatem poniższe wymogi zawsze byłyby stosowane, biorąc pod uwagę znaczenie tego zastosowania dla osób fizycznych oraz w świetle dorobku prawnego UE w dziedzinie równości zatrudnienia. Możliwe byłoby rozważenie dalszych konkretnych zastosowań mających wpływ na prawa konsumentów;
- zastosowanie AI do celów zdalnej identyfikacji biometrycznej⁵² i innych inwazyjnych technologii nadzoru zawsze byłoby uznawane za charakteryzujące się wysokim ryzykiem, a zatem poniższe wymogi byłyby zawsze stosowane.

D. RODZAJE WYMOGÓW

Przy opracowywaniu przyszłych ram regulacyjnych dotyczących sztucznej inteligencji konieczne będzie podjęcie decyzji w sprawie rodzajów obowiązkowych wymogów prawnych nakładanych na odpowiednie podmioty. Wymogi te mogą zostać bardziej szczegółowo określone w oparciu o normy. Jak zauważono w sekcji C powyżej oraz oprócz już istniejącego prawodawstwa, wymogi te miałyby zastosowanie wyłącznie do zastosowań AI wysokiego ryzyka, zapewniając w ten sposób ukierunkowanie i proporcjonalność wszelkiej interwencji regulacyjnej.

Biorąc pod uwagę wytyczne grupy ekspertów wysokiego szczebla oraz kwestie przedstawione powyżej, wymogi dotyczące zastosowań sztucznej inteligencji charakteryzujących się wysokim ryzykiem mogłyby składać się z następujących kluczowych elementów, które omówiono bardziej szczegółowo w podsekcjach poniżej:

- dane szkoleniowe;
- przechowywanie danych i prowadzenie rejestrów;

⁵¹ Należy zaznaczyć, że zastosowanie mogą mieć również inne akty prawne UE. Na przykład dyrektywa w sprawie ogólnego bezpieczeństwa produktów może mieć zastosowanie do bezpieczeństwa zastosowań AI, jeżeli są one wbudowane w produkt konsumpcyjny.

⁵² Należy odróżnić zdalną identyfikację biometryczną od uwierzytelniania biometrycznego (czyli procedury bezpieczeństwa, która opiera się na niepowtarzalnej charakterystyce biologicznej osoby w celu sprawdzenia, czy osoba jest tym, za kogo się podaje). Zdalna identyfikacja biometryczna to proces, w ramach którego tożsamość wielu osób jest ustalana z pomocą identyfikatorów biometrycznych (odcisków palców, wizerunku twarzy, tęczy, tęczówek, układu naczyń krwionośnych itp.) na odległość, w przestrzeni publicznej oraz w sposób ciągły lub trwały, poprzez porównywanie z danymi przechowywanymi w bazie danych.

- wymagane informacje;
- solidność i dokładność;
- sprawowanie nadzoru przez człowieka;
- szczególne wymogi dotyczące określonych zastosowań sztucznej inteligencji, np. do celów zdalnej identyfikacji biometrycznej.

Aby zapewnić pewność prawa, wymogi te zostaną doprecyzowane i staną się jasnym punktem odniesienia dla wszystkich podmiotów, które muszą je spełnić.

a) Dane szkoleniowe

Promowanie, wzmacnianie i obrona wartości i zasad UE, a w szczególności praw, które przysługują obywatelom na mocy prawodawstwa UE, są ważniejsze niż kiedykolwiek dotychczas. Wysiłki te obejmują również omawiane w niniejszym dokumencie zastosowania AI wysokiego ryzyka, które są wprowadzane do obrotu i wykorzystywane w UE.

Jak już wspomniano, bez danych nie ma sztucznej inteligencji. Funkcjonowanie wielu systemów sztucznej inteligencji oraz podejmowane w rezultacie działania i decyzje w dużym stopniu zależą od zbioru danych, na którym systemy zostały przeszkolone. Należy zatem wprowadzić niezbędne środki w celu zapewnienia, aby w przypadku danych wykorzystywanych do szkolenia systemów sztucznej inteligencji przestrzegane były wartości i zasady UE, w szczególności w odniesieniu do bezpieczeństwa, oraz obowiązujące przepisy prawa dotyczące ochrony praw podstawowych. Można przewidzieć następujące wymogi dotyczące zbiorów danych wykorzystywanych do szkolenia systemów sztucznej inteligencji:

- wymogi mające na celu uzyskanie wystarczającej pewności, że późniejsze wykorzystanie produktów lub usług, które umożliwia zastosowanie systemu sztucznej inteligencji, jest bezpieczne, ponieważ spełnia normy określone w stosownych unijnych przepisach bezpieczeństwa (istniejących i ewentualnych uzupełniających). Na przykład wymogi zapewniające szkolenie systemów sztucznej inteligencji na zbiorach danych, które są wystarczająco duże i obejmują wszystkie odpowiednie scenariusze niezbędne do uniknięcia niebezpiecznych sytuacji;
- wymogi dotyczące stosowania racjonalnych środków mających na celu zapewnienie, by w wyniku późniejszego stosowania systemów AI nie dochodziło do wystąpienia zakazanej dyskryminacji. Wymogi te mogą obejmować w szczególności obowiązek stosowania zbiorów danych, które są wystarczająco reprezentatywne, zwłaszcza w celu zapewnienia odpowiedniego uwzględnienia w nich wszystkich istotnych aspektów płci, pochodzenia etnicznego i innych możliwych powodów zakazanej dyskryminacji;
- wymogi mające na celu zapewnienie odpowiedniej ochrony prywatności i danych osobowych podczas korzystania z produktów i usług opartych na sztucznej inteligencji. Ogólne rozporządzenie o ochronie danych i dyrektywa w sprawie egzekwowania prawa regulują te kwestie w odniesieniu do problematyki wchodzącej w ich zakres.

b) Prowadzenie rejestrów i przechowywanie danych

Biorąc pod uwagę elementy takie jak złożoność i nieprzejrzystość wielu systemów sztucznej inteligencji i związane z tym trudności, które mogą zaistnieć w odniesieniu do skutecznego sprawdzania zgodności z obowiązującymi przepisami i ich egzekwowania, konieczne jest wprowadzenie wymogów dotyczących prowadzenia rejestrów w odniesieniu do programowania algorytmu, danych wykorzystywanych do szkolenia systemów sztucznej inteligencji wysokiego

ryzyka oraz, w niektórych przypadkach, przechowywania samych danych. Wymogi te zasadniczo pozwalają na prześledzenie i weryfikację potencjalnie problematycznych działań lub decyzji podejmowanych przez systemy sztucznej inteligencji. Powinno to nie tylko ułatwić nadzór i egzekwowanie przepisów, ale może również stanowić zachętę dla zainteresowanych podmiotów gospodarczych do uwzględnienia konieczności przestrzegania tych przepisów na wczesnym etapie.

W tym celu ramy regulacyjne mogłyby przewidywać przechowywanie:

- dokładnych rejestrów dotyczących zbioru danych wykorzystywanych do szkolenia i testowania systemów sztucznej inteligencji, w tym opisu głównych cech i sposobu wyboru zbioru danych;
- w niektórych uzasadnionych przypadkach, samych zbiorów danych;
- dokumentacji dotyczącej metod, procesów i technik programowania⁵³ i szkolenia stosowanych do budowy, testowania i walidacji systemów sztucznej inteligencji, w tym, w stosownych przypadkach, w odniesieniu do bezpieczeństwa i unikania uprzedzeń, które mogłyby prowadzić do zakazanej dyskryminacji.

Rejestry, dokumentacja i w stosownych przypadkach zbiory danych, powinny być przechowywane przez ograniczony, rozsądny okres w celu zapewnienia skutecznego egzekwowania odpowiednich przepisów. Należy wprowadzić środki w celu zapewnienia udostępniania ich na żądanie, w szczególności na potrzeby testów lub kontroli przeprowadzanych przez właściwe organy. W razie konieczności należy dokonać ustaleń w celu zapewnienia ochrony informacji poufnych, takich jak tajemnice handlowe.

c) Dostarczanie informacji

Wymóg przejrzystości wykracza poza wymogi w zakresie prowadzenia rejestrów, o których mowa w lit. c) powyżej. Aby osiągnąć wyznaczone cele – w szczególności promowanie odpowiedzialnego korzystania ze sztucznej inteligencji, budowanie zaufania i ułatwienie dochodzenia roszczeń w razie potrzeby – ważne jest, aby w sposób proaktywny przekazywane były odpowiednie informacje na temat stosowania systemów AI wysokiego ryzyka.

W związku z tym można rozważyć następujące wymogi:

- zapewnienie przekazywania jasnych informacji na temat możliwości i ograniczeń systemu AI, w szczególności celu, do którego dany system jest przeznaczony, warunków funkcjonowania zgodnie z przeznaczeniem oraz oczekiwanego poziomu dokładności w osiąganiu określonego celu. Informacje te są ważne zwłaszcza dla operatorów systemów, ale mogą również być istotne dla właściwych organów i stron, na które system ma wpływ;
- ponadto obywatele powinni być wyraźnie informowani, że kontaktują się z systemem sztucznej inteligencji, a nie z człowiekiem. Choć prawodawstwo UE w zakresie ochrony danych zawiera już pewne przepisy tego rodzaju⁵⁴, można by wprowadzić dodatkowe wymogi służące osiągnięciu wyżej wymienionych celów. Należy unikać niepotrzebnych obciążeń.

⁵³ Na przykład dokumentacja dotycząca algorytmu, w tym informacje o tym, do jakich celów model jest zoptymalizowany, jakie wagi zaprojektowano w odniesieniu do określonych parametrów na wejściu itd.

⁵⁴ W szczególności zgodnie z art. 13 ust. 2 lit. f) RODO administratorzy muszą, podczas pozyskiwania danych osobowych, podać osobom, których dane dotyczą, dodatkowe informacje niezbędne do zapewnienia rzetelności i przejrzystości przetwarzania, dotyczące zautomatyzowanego podejmowania decyzji, jak również określone inne informacje.

W związku z tym nie ma konieczności udzielania takich informacji na przykład w sytuacjach, w których obywatele od razu wiedzą, że kontaktują się z systemami sztucznej inteligencji. Ponadto istotne jest, aby dostarczone informacje były obiektywne, zwarte i łatwe do zrozumienia. Sposób dostarczania informacji powinien być dostosowany do danego kontekstu.

d) Solidność i dokładność

Aby systemy sztucznej inteligencji – a w szczególności zastosowania AI wysokiego ryzyka – mogły być wiarygodne, muszą być solidne pod względem technicznym i dokładne. Oznacza to, że takie systemy muszą być opracowywane w odpowiedzialny sposób i z należyтым uprzednim uwzględnieniem ryzyka, jakie mogą generować. Ich rozwój i funkcjonowanie muszą zapewniać wiarygodne działanie zgodnie z przeznaczeniem. Należy wprowadzić wszelkie racjonalne środki w celu zminimalizowania ryzyka wystąpienia szkody.

W związku z tym można rozważyć następujące elementy:

- wymogi zapewniające, aby systemy sztucznej inteligencji były solidne i dokładne lub przynajmniej wymóg prawidłowego odzwierciedlenia poziomu dokładności, na wszystkich etapach cyklu życia;
- wymogi zapewniające odtwarzalność wyników;
- wymogi zapewniające, aby systemy sztucznej inteligencji mogły w odpowiedni sposób radzić sobie z błędami lub niespójnościami w całym cyklu życia;
- wymogi zapewniające, aby systemy sztucznej inteligencji były odporne zarówno na jawne ataki, jak i na bardziej subtelne próby manipulacji danymi lub algorytmami, oraz aby w takich przypadkach podejmowane były środki łagodzące.

e) Sprawowanie nadzoru przez człowieka

Sprawowanie nadzoru przez człowieka pomaga zapewnić, aby system AI nie podważał autonomii człowieka ani nie wywierał innych niekorzystnych skutków. Cel, jakim jest wiarygodna, etyczna i ukierunkowana na człowieka sztuczna inteligencja, można osiągnąć jedynie poprzez zapewnienie odpowiedniego zaangażowania człowieka w odniesieniu do zastosowań AI wysokiego ryzyka.

Chociaż wszystkie zastosowania AI rozważane w niniejszej białej księdze w odniesieniu do konkretnego systemu prawnego uznaje się za charakteryzujące się wysokim ryzykiem, to właściwy rodzaj i stopień nadzoru przez człowieka mogą różnić się w poszczególnych przypadkach. Zależy on w szczególności od zamierzonego zastosowania systemów oraz od skutków, jakie zastosowanie mogłoby mieć dla zainteresowanych obywateli i osób prawnych. Nie może on również naruszać praw ustanowionych w RODO, jeżeli system sztucznej inteligencji przetwarza dane osobowe. Na przykład nadzór przez człowieka mógłby przejawiać się w następujący sposób:

- wynik działania systemu sztucznej inteligencji nie wywiera skutku, jeżeli nie został wcześniej poddany kontroli i zatwierdzony przez człowieka (np. odrzucenie wniosku o świadczenia z zabezpieczenia społecznego może być dokonane wyłącznie przez człowieka);
- wynik działania systemu sztucznej inteligencji wywiera natychmiastowy skutek, ale następnie zapewniona jest interwencja człowieka (np. odrzucenie wniosku o kartę kredytową może

zostać przeprowadzone przez system sztucznej inteligencji, ale później musi być możliwa kontrola tej decyzji przez człowieka);

- monitorowanie systemu sztucznej inteligencji w trakcie jego eksploatacji oraz zdolność do interweniowania w czasie rzeczywistym i dezaktywacji (np. przycisk lub procedura zatrzymania dostępne w samochodzie autonomicznym, jeżeli człowiek stwierdzi, że działanie samochodu nie jest bezpieczne);
- na etapie projektowania, poprzez nałożenie ograniczeń operacyjnych na system sztucznej inteligencji (np. samochód autonomiczny przestaje działać w warunkach słabej widoczności, kiedy czujniki mogą być mniej wiarygodne, lub niezależnie od warunków zachowuje określoną odległość od pojazdu znajdującego się przed nim).

f) Szczegółowe wymogi dotyczące zdalnej identyfikacji biometrycznej

Gromadzenie i wykorzystywanie danych biometrycznych⁵⁵ do celów zdalnej identyfikacji⁵⁶, na przykład poprzez wykorzystanie rozpoznawania twarzy w miejscach publicznych, niesie ze sobą szczególne ryzyko dla praw podstawowych⁵⁷. Wpływ korzystania z systemów AI służących zdalnej identyfikacji biometrycznej na prawa podstawowe może się znacznie różnić w zależności od celu, kontekstu i zakresu zastosowania.

Unijne przepisy dotyczące ochrony danych zakazują co do zasady przetwarzania danych biometrycznych w celu jednoznacznego zidentyfikowania osoby fizycznej, z wyjątkiem określonych warunków⁵⁸. W szczególności na mocy RODO przetwarzanie takie może odbywać się tylko z ograniczonej liczby powodów, z których głównym jest istotny interes publiczny. W takim przypadku przetwarzanie musi odbywać się na podstawie prawa unijnego lub krajowego, z zastrzeżeniem wymogów proporcjonalności, poszanowania istoty prawa do ochrony danych i odpowiednich zabezpieczeń. Zgodnie z dyrektywą w sprawie egzekwowania prawa musi istnieć bezwzględna potrzeba takiego przetwarzania, co do zasady dopuszczenie na mocy prawa unijnego lub krajowego, a także odpowiednie zabezpieczenia. Jako że każdy przypadek przetwarzania danych biometrycznych w celu jednoznacznej identyfikacji osoby fizycznej stanowiłby wyjątek od zakazu ustanowionego w prawie UE, będzie on podlegał zasadom Karty praw podstawowych Unii Europejskiej.

⁵⁵ Dane biometryczne definiuje się jako dane osobowe, które wynikają ze specjalnego przetwarzania technicznego, dotyczą cech fizycznych, fizjologicznych lub behawioralnych osoby fizycznej oraz umożliwiają lub potwierdzają jednoznaczne uwierzytelnienie lub identyfikację tej osoby, takie jak wizerunek twarzy lub dane daktyloskopijne [odciski palców] (dyrektywa w sprawie egzekwowania prawa, art. 3 pkt 13; RODO, art. 4 pkt 14, rozporządzenie (UE) 2018/1725, art. 3 pkt 18).

⁵⁶ W odniesieniu do rozpoznawania twarzy identyfikacja oznacza, że wzór obrazu twarzy danej osoby jest porównywany z wieloma innymi wzorami przechowywanymi w bazie danych w celu stwierdzenia, czy obraz twarzy tej osoby jest przechowywany w tej bazie danych. Uwierzytelnienie (lub weryfikacja) natomiast często odnosi się do porównania „jeden do jednego”. Umożliwia ono porównanie dwóch wzorców biometrycznych, co do których zasadniczo zakłada się, że należą do tej samej osoby. Dwa wzorce biometryczne są porównywane w celu określenia, czy osoba wskazana na dwóch obrazach jest tą samą osobą. Taka procedura jest na przykład stosowana w bramkach zautomatyzowanej kontroli granicznej wykorzystywanych do odprawy granicznej w portach lotniczych.

⁵⁷ Na przykład w odniesieniu do godności ludzkiej. Jeśli chodzi o korzystanie z technologii rozpoznawania twarzy, prawo do poszanowania życia prywatnego i ochrony danych osobowych jest ważnym elementem kwestii związanych z prawami podstawowymi. Występuje również potencjalny wpływ związany z brakiem dyskryminacji i prawami grup specjalnych, takich jak dzieci, osoby starsze i osoby z niepełnosprawnościami. Ponadto stosowanie tej technologii nie może ograniczać wolności słowa, zrzeszania się i zgromadzeń. Zob.: Facial recognition technology: fundamental rights considerations in the context of law enforcement (Technologia rozpoznawania twarzy: kwestie związane z prawami podstawowymi w kontekście egzekwowania prawa), <https://fra.europa.eu/en/publication/2019/facial-recognition>.

⁵⁸ Art. 9 RODO, art. 10 dyrektywy w sprawie egzekwowania prawa. Zob. również art. 10 rozporządzenia (UE) 2018/1725 (mającego zastosowanie do instytucji i organów UE).

Z powyższego wynika, że zgodnie z obowiązującymi unijnymi przepisami o ochronie danych i Kartą praw podstawowych sztuczna inteligencja może być wykorzystywana do celów zdalnej identyfikacji biometrycznej tylko wtedy, gdy takie wykorzystanie jest należycie uzasadnione, proporcjonalne i podlega odpowiednim zabezpieczeniom.

W celu uwzględnienia możliwych problemów społecznych związanych ze stosowaniem sztucznej inteligencji do takich celów w miejscach publicznych oraz w celu uniknięcia rozdrobnienia rynku wewnętrznego, Komisja rozpocznie szeroko zakrojoną debatę europejską na temat konkretnych okoliczności (jeżeli takie istnieją), które mogłyby uzasadnić takie zastosowanie, oraz na temat wspólnych zabezpieczeń.

E. ADRESACI

W odniesieniu do adresatów wymogów prawnych, które obowiązywałyby w odniesieniu do wyżej omówionych zastosowań AI wysokiego ryzyka, należy rozważyć dwie główne kwestie.

Po pierwsze pojawia się pytanie, w jaki sposób należy rozdzielić obowiązki między zainteresowane podmioty gospodarcze. W cykl życia systemu sztucznej inteligencji zaangażowanych jest wiele podmiotów. Należą do nich deweloper, operator (osoba korzystająca z produktu lub usługi opartych na AI) i potencjalnie inne podmioty (producent, dystrybutor lub importer, dostawca usług, użytkownik profesjonalny lub prywatny).

Komisja jest zdania, że w przyszłych ramach regulacyjnych każdy obowiązek powinien być skierowany do podmiotu lub podmiotów, które najlepiej mogą zająć się wszelkimi potencjalnymi zagrożeniami. Na przykład, chociaż deweloperzy sztucznej inteligencji mogą najlepiej zająć się zagrożeniami wynikającymi z fazy rozwoju, ich zdolność do kontrolowania ryzyka w fazie użytkowania może być bardziej ograniczona. W takim przypadku odpowiedni obowiązek należałoby nałożyć na operatora. Pozostaje to bez uszczerbku dla kwestii, która strona powinna ponosić odpowiedzialność za wszelkie wyrządzone szkody w odniesieniu do odpowiedzialności wobec użytkowników końcowych lub innych stron, które poniosły szkodę i zapewnienia skutecznego dostępu do wymiaru sprawiedliwości. Zgodnie z unijnym prawem dotyczącym odpowiedzialności za produkt, odpowiedzialność za produkty wadliwe przypisuje się producentowi, bez uszczerbku dla przepisów krajowych, które mogą również umożliwić dochodzenie roszczeń od innych stron.

Po drugie pojawia się pytanie o zasięg geograficzny interwencji legislacyjnej. Zdaniem Komisji nadrzędne znaczenie ma to, by wymogi miały zastosowanie do wszystkich odpowiednich podmiotów gospodarczych, które dostarczają produkty lub świadczą usługi oparte na AI na terytorium Unii, niezależnie od tego, czy mają w Unii siedzibę. W przeciwnym razie wspomniane wcześniej cele interwencji legislacyjnej mogłyby nie zostać w pełni zrealizowane.

F. PRZESTRZEGANIE I EGZEKOWANIE PRAWA

W celu zapewnienia, by sztuczna inteligencja była godna zaufania i bezpieczna oraz przestrzegała europejskich wartości i zasad, obowiązujące wymogi prawne muszą być przestrzegane w praktyce i być skutecznie egzekwowane zarówno przez właściwe organy krajowe i europejskie, jak i przez zainteresowane strony. Właściwe organy powinny mieć możliwość badania indywidualnych przypadków, ale również dokonania oceny wpływu na społeczeństwo.

Z uwagi na wysokie ryzyko, jakie niektóre zastosowania AI stwarzają dla obywateli i naszego społeczeństwa (zob. sekcja A powyżej), na obecnym etapie Komisja uważa, że konieczna byłaby obiektywna i uprzednia ocena zgodności w celu sprawdzenia i zapewnienia przestrzegania niektórych z wyżej wymienionych obowiązkowych wymogów dotyczących zastosowań wysokiego ryzyka (zob.

sekcja D powyżej). Uprzednia ocena zgodności może obejmować procedury testowania, kontroli lub certyfikacji⁵⁹. Może ona obejmować kontrole algorytmów i zbiorów danych wykorzystywanych w fazie rozwoju.

Oceny zgodności dotyczące zastosowań AI wysokiego ryzyka powinny stanowić część mechanizmów oceny zgodności, które już istnieją w odniesieniu do dużej liczby produktów wprowadzanych na rynek wewnętrzny UE. Jeżeli nie można polegać na już istniejących mechanizmach, może zaistnieć potrzeba ustanowienia podobnych mechanizmów w oparciu o najlepsze praktyki i ewentualny wkład zainteresowanych stron i europejskich organizacji normalizacyjnych. Każdy nowy mechanizm powinien być proporcjonalny i niedyskryminacyjny oraz być stosowany w oparciu o przejrzyste i obiektywne kryteria zgodności z zobowiązaniami międzynarodowymi.

Przy opracowywaniu i wdrażaniu systemu opartego na uprzednich ocenach zgodności szczególną uwagę należy zwrócić na następujące kwestie:

- nie wszystkie wymogi przedstawione powyżej mogą być odpowiednie do weryfikacji w ramach uprzedniej oceny zgodności. Na przykład wymóg dotyczący informacji, które należy dostarczyć, na ogół nie nadaje się dobrze do weryfikacji w drodze takiej oceny.
- Szczególną uwagę należy zwrócić na możliwość ewolucji niektórych systemów sztucznej inteligencji i uczenia się na podstawie doświadczeń, co może wymagać powtarzania ocen w ciągu całego cyklu życia tych systemów;
- potrzeba weryfikacji danych wykorzystywanych do szkolenia oraz odpowiednich metod, procesów i technik programowania i szkolenia stosowanych do budowy, testowania i walidacji systemów sztucznej inteligencji;
- jeżeli z oceny zgodności wynika, że system sztucznej inteligencji nie spełnia wymogów, na przykład dotyczących danych wykorzystywanych do szkolenia, stwierdzone niedociągnięcia będą musiały zostać usunięte, na przykład poprzez ponowne przeszkolenie systemu w UE w taki sposób, aby zapewnić spełnienie wszystkich mających zastosowanie wymogów.

Oceny zgodności byłyby obowiązkowe dla wszystkich podmiotów gospodarczych, do których odnoszą się wymogi, niezależnie od siedziby⁶⁰. Aby ograniczyć obciążenie dla MŚP, można by rozważyć strukturę wsparcia, w tym za pośrednictwem ośrodków innowacji cyfrowych. Ponadto normy oraz specjalne narzędzia internetowe mogłyby ułatwić przestrzeganie przepisów.

Wszelkie uprzednie oceny zgodności powinny pozostawać bez uszczerbku dla monitorowania zgodności z przepisami i egzekwowania przepisów *ex post* przez właściwe organy krajowe. Odnosi się to do zastosowań AI wysokiego ryzyka, ale również do innych zastosowań w dziedzinie sztucznej inteligencji podlegających wymogom prawnym, chociaż ze względu na wysokie ryzyko przedmiotowych zastosowań właściwe organy krajowe mogłyby zwrócić szczególną uwagę właśnie na nie. N potrzeby kontroli *ex post* powinna być dostępna odpowiednia dokumentacja dotycząca przedmiotowych zastosowań AI (zob. sekcja E powyżej) oraz, w stosownych przypadkach, powinna

⁵⁹ System opierałby się na procedurach oceny zgodności w UE, zob. decyzja 768/2008/WE lub rozporządzenie (UE) 2019/881 (akt o cyberbezpieczeństwie), z uwzględnieniem specyfiki sztucznej inteligencji. Zob. Niebieski przewodnik – wdrażanie unijnych przepisów dotyczących produktów, 2014 r.

⁶⁰ W odniesieniu do właściwej struktury zarządzania, w tym organów wyznaczonych do przeprowadzania ocen zgodności – zob. sekcja H poniżej.

istnieć możliwość testowania takich zastosowań przez osoby trzecie, takie jak właściwe organy. Może to być szczególnie ważne, gdy zachodzi ryzyko związane z prawami podstawowymi, które zależą od kontekstu. Takie monitorowanie zgodności powinno być częścią ciągłego systemu nadzoru rynku. Aspekty związane z zarządzaniem omówiono bardziej szczegółowo w sekcji H poniżej.

Ponadto należy zapewnić skuteczne dochodzenie odszkodowania na drodze sądowej dla stron dotkniętych negatywnymi skutkami systemów sztucznej inteligencji zarówno w odniesieniu do zastosowań AI wysokiego ryzyka, jak i innych. Kwestie odpowiedzialności są przedmiotem dalszej dyskusji w sprawozdaniu w sprawie ram bezpieczeństwa i odpowiedzialności towarzyszącym niniejszej białej księdze.

G. DOBROWOLNE ETYKIETOWANIE ZASTOSOWAŃ AI NIECHARAKTERYZUJĄCYCH SIĘ WYSOKIM RYZYKIEM

W przypadku zastosowań AI, które nie kwalifikują się jako zastosowania „wysokiego ryzyka” (zob. sekcja C powyżej) i w związku z tym nie podlegają obowiązkowym wymogom omówionym powyżej (zob. sekcje D, E i F), oprócz mającego zastosowanie ustawodawstwa istnieje możliwość ustanowienia systemu dobrowolnego etykietowania.

W ramach tego systemu zainteresowane podmioty gospodarcze, które nie są objęte obowiązkowymi wymogami, mogą podjąć decyzję o dobrowolnym poddaniu się tym wymogom albo określone mu zbiorowi podobnych wymogów ustanowionych specjalnie na potrzeby dobrowolnego systemu. Zainteresowane podmioty gospodarcze uzyskałyby wówczas znak jakości dla swoich zastosowań sztucznej inteligencji.

Dobrowolna etykieta umożliwiłaby zainteresowanym podmiotom gospodarczym sygnalizowanie, że ich produkty i usługi oparte na sztucznej inteligencji są godne zaufania. Dzięki temu użytkownicy mogliby z łatwością rozpoznać, że dane produkty i usługi są zgodne z określonymi obiektywnymi i znormalizowanymi ogólnounijnymi wskaźnikami, które wykraczają poza zazwyczaj obowiązujące zobowiązania prawne. Pomogłoby to zwiększyć zaufanie użytkowników do systemów sztucznej inteligencji i promować ogólne wykorzystanie tej technologii.

Wariant ten pociągałby za sobą utworzenie nowego instrumentu prawnego, który określałby ramy dobrowolnego etykietowania dla deweloperów i/lub operatorów systemów sztucznej inteligencji, które nie są uznawane za charakteryzujące się wysokim ryzykiem. Uczestnictwo w systemie etykietowania byłoby dobrowolne, ale w momencie gdyby deweloper lub operator zdecydował się na stosowanie etykiety, wymogi stałyby się wiążące. Połączenie egzekwowania *ex ante* i *ex post* musiałyby zapewnić spełnienie wszystkich wymogów.

H. ZARZĄDZANIE

Konieczne jest utworzenie europejskiej struktury zarządzania w zakresie sztucznej inteligencji w formie ram współpracy właściwych organów krajowych, aby uniknąć rozdrobnienia odpowiedzialności, zwiększyć zdolności w państwach członkowskich i zapewnić, by Europa stopniowo uzyskiwała zdolności niezbędne do testowania i certyfikacji produktów i usług opartych na sztucznej inteligencji. W tym kontekście korzystne byłoby wspieranie właściwych organów krajowych, aby mogły one wypełniać swój mandat w odniesieniu do stosowania sztucznej inteligencji.

Europejska struktura zarządzania mogłaby mieć różne zadania – stanowić forum regularnej wymiany informacji i najlepszych praktyk, określać pojawiające się tendencje czy też doradzać w zakresie działalności normalizacyjnej i certyfikacji. Powinna również odgrywać kluczową rolę w ułatwianiu wdrażania ram prawnych, np. poprzez wydawanie wytycznych i opinii lub zapewnianie wiedzy

fachowej. W związku z tym powinna opierać się na sieci organów krajowych, a także na sieciach sektorowych i organach regulacyjnych na szczeblu krajowym i unijnym. Ponadto wsparcia Komisji mógłby udzielać komitet ekspertów.

Struktura zarządzania powinna gwarantować maksymalne uczestnictwo zainteresowanych stron. Z zainteresowanymi stronami – organizacjami konsumentów i partnerami społecznymi, przedsiębiorstwami, naukowcami i organizacjami społeczeństwa obywatelskiego – należy prowadzić konsultacje na temat wdrażania i dalszego rozwoju tych ram.

Biorąc pod uwagę istniejące już struktury w sektorach finansów, produktów leczniczych, lotnictwa, wyrobów medycznych, ochrony konsumentów i ochrony danych, proponowana struktura zarządzania nie powinna powielać istniejących funkcji. Zamiast tego należy ustanowić ścisłe powiązania z innymi unijnymi i krajowymi właściwymi organami w różnych sektorach, aby uzupełnić istniejącą wiedzę fachową i pomóc istniejącym organom w monitorowaniu i nadzorowaniu działalności podmiotów gospodarczych stosujących systemy sztucznej inteligencji i produkty lub usługi oparte na sztucznej inteligencji.

Jeżeli wariant ten zostanie wybrany, przeprowadzanie ocen zgodności może być powierzone jednostkom notyfikowanym wyznaczonym przez państwa członkowskie. Ośrodki badawcze powinny umożliwiać niezależny audyt i ocenę systemów AI zgodnie z wymogami określonymi powyżej. Niezależna ocena zwiększy zaufanie i zapewni obiektywność. Mogłaby również ułatwić pracę odpowiednich właściwych organów.

UE dysponuje doskonałymi ośrodkami badań i oceny, przy czym powinna rozwijać swoje zdolności również w dziedzinie sztucznej inteligencji. Podmioty gospodarcze mające siedzibę w państwach trzecich, które chcą wejść na rynek wewnętrzny, mogłyby korzystać z usług wyznaczonych organów mających siedzibę w UE lub, z zastrzeżeniem umów o wzajemnym uznaniu z państwami trzecimi, korzystać z usług organów państw trzecich wyznaczonych do przeprowadzenia takiej oceny.

Struktura zarządzania dotycząca sztucznej inteligencji i ewentualne oceny zgodności nie miałyby wpływu na wynikające z obowiązującego prawa UE uprawnienia i zakres odpowiedzialności odpowiednich właściwych organów w określonych sektorach lub w określonych kwestiach (finanse, produkty lecznicze, lotnictwo, wyroby medyczne, ochrona konsumentów, ochrona danych itp.).

6. PODSUMOWANIE

Sztuczna inteligencja jest strategiczną technologią, która przynosi wiele korzyści obywatelom, przedsiębiorstwom i całemu społeczeństwu, pod warunkiem że jest ukierunkowana na człowieka, etyczna, zrównoważona i przestrzega podstawowych praw i wartości. AI oferuje istotne korzyści związane z efektywnością i wydajnością, które mogą wzmocnić konkurencyjność przemysłu europejskiego i poprawić dobrostan obywateli. Może również przyczynić się do znalezienia rozwiązań niektórych z najpilniejszych wyzwań społecznych, w tym związanych z przeciwdziałaniem zmianie klimatu i degradacją środowiska, wyzwań związanych ze zrównoważonym rozwojem i zmianami demograficznymi oraz ochroną demokracji, a także – w razie potrzeby i w sposób proporcjonalny – przyczynić się do walki z przestępczością.

Aby Europa mogła w pełni wykorzystać możliwości oferowane przez AI, musi rozwijać i wzmacniać niezbędne zdolności przemysłowe i technologiczne. Jak określono w europejskiej strategii w zakresie danych towarzyszącej niniejszej białej księdze, wymaga to również środków, które pozwolą UE stać się globalnym centrum danych.

Europejskie podejście do sztucznej inteligencji ma na celu promowanie potencjału innowacyjnego Europy w tej dziedzinie, przy jednoczesnym wspieraniu rozwoju i wprowadzania etycznej i godnej zaufania sztucznej inteligencji w całą gospodarkę unijną. Sztuczna inteligencja powinna działać na rzecz ludzi i społeczeństwa.

W oparciu o niniejszą białą księgę oraz towarzyszące jej sprawozdanie w sprawie ram bezpieczeństwa i odpowiedzialności Komisja rozpoczyna szeroko zakrojone konsultacje z przedstawicielami społeczeństwa obywatelskiego, przemysłu i środowiska akademickiego państw członkowskich dotyczące konkretnych propozycji europejskiego podejścia do sztucznej inteligencji. Obejmują one zarówno środki polityczne mające na celu pobudzenie inwestycji w badania i innowacje, zwiększenie rozwoju umiejętności i wspieranie wykorzystania sztucznej inteligencji przez MŚP, jak i propozycje dotyczące kluczowych elementów przyszłych ram regulacyjnych. Konsultacje te umożliwią prowadzenie kompleksowego dialogu ze wszystkimi zainteresowanymi stronami, którego wyniki przyczynią się do decyzji o kolejnych krokach Komisji.

Komisja zachęca do zgłaszania uwag na temat propozycji przedstawionych w białej księdze poprzez otwarte konsultacje publiczne dostępne na stronie internetowej https://ec.europa.eu/info/consultations_pl. Konsultacje są otwarte do dnia 19 maja 2020 r.

Standardową praktyką Komisji jest publikowanie uwag otrzymanych w odpowiedzi na konsultacje publiczne. Można jednak zwrócić się o poufne traktowanie całości lub części nadesłanych uwag. W takim wypadku należy wyraźnie zaznaczyć na stronie tytułowej dokumentu zawierającego uwagi, że nie należy ich publikować, a ponadto przesłać Komisji wersję przeznaczoną do publikacji, pozbawioną elementów poufnych.