

Opis projektu badawczego nr 2021/41/B/HS5/01001

„Zastosowanie autonomicznej AI w procesie pracy a odpowiedzialność cywilna pracodawcy wobec osób trzecich”

Kierownik projektu: dr Iwona Gredka-Ligarska

1) Cel badawczy projektu. Celem niniejszego projektu badawczego jest ustalenie, jaki model odpowiedzialności cywilnej powinien znaleźć zastosowanie w prawodawstwie Unii Europejskiej w sytuacji, gdy autonomiczna sztuczna inteligencja (AI – artificial intelligence), wykorzystywana przez pracodawcę w procesie pracy, wyrządza szkodę osobie trzeciej. Wybrany cel badawczy ma charakter nowatorski i powinien zostać zrealizowany z uwagi na rewolucyjne zmiany wywołane zastosowaniem autonomicznej sztucznej inteligencji. Obecnie odpowiedzialność cywilna za szkodę opiera się na utrwalonym podziale na osoby i rzeczy. W konsekwencji producent ponosi odpowiedzialność cywilną za szkodę wyrządzoną przez produkt niebezpieczny (rzecz), natomiast pracodawca ponosi odpowiedzialność cywilną za szkodę wyrządzoną osobie trzeciej przez pracownika (człowieka) w związku z wykonywaniem pracy. Sztuczna inteligencja nie jest rzeczą i nigdy nie stanie się człowiekiem. Istnieje jednak wiele przesłanek wskazujących, że w ramach procesu pracy AI będzie zastępować nie tylko rzeczy (maszyny, urządzenia), lecz także ludzi (pracowników). Oczywiście proces ten nie nastąpi natychmiast ani nie obejmie wszystkich zawodów. Niemniej jednak w chwili obecnej podejmowane są badania naukowe zmierzające do ustalenia, które profesje zostaną w przyszłości wyparte z rynku przez autonomiczną AI (Gentili, Compagnucci, Gallegati, Valentini, 2020). W tym kontekście dotychczasowe dostosowanie przesłanek i modeli odpowiedzialności cywilnej do rozróżnienia między osobami a rzeczami okazuje się niewystarczające. W prawodawstwie Unii Europejskiej prawodawca nie zdecydował się dotychczas na przyznanie sztucznej inteligencji podmiotowości prawnej, a tym samym na ustanowienie jej samodzielnej odpowiedzialności cywilnej za szkodę. Nie można jednak wykluczyć konieczności uprzedniego wypracowania całkowicie nowego, dotychczas nieznanego modelu odpowiedzialności odnoszącego się do autonomicznej AI. Być może zasadne okaże się również zastąpienie różnych reżimów odpowiedzialności cywilnej (odpowiedzialności pracodawcy wobec osób trzecich za szkody wyrządzone przez pracowników; odpowiedzialności prowadzącego przedsiębiorstwo lub zakład wprawiany w ruch za pomocą sił przyrody; odpowiedzialności producenta za maszyny i urządzenia wykorzystywane w procesie pracy) jednym, jednolitym typem odpowiedzialności dedykowanym autonomicznej sztucznej inteligencji. Hipoteza badawcza opiera się na założeniu, że reżim odpowiedzialności powinien zostać dostosowany do funkcjonowania w pełni autonomicznej AI, która w określonym obszarze będzie podejmować samodzielne decyzje bez udziału człowieka. Szczegółowym celem badawczym jest udzielenie odpowiedzi na pytanie: kto powinien ponosić odpowiedzialność za szkodę wyrządzoną osobie trzeciej przez autonomiczną AI wykorzystywaną przez pracodawcę w procesie pracy oraz na jakich zasadach odpowiedzialność ta powinna się opierać? Zgodnie z projektem Rozporządzenia Parlamentu Europejskiego i Rady ustanawiającego zharmonizowane przepisy dotyczące sztucznej inteligencji (Artificial Intelligence Act) oraz zmieniającego niektóre akty ustawodawcze Unii (Bruksela, 21.4.2021 r.), „system sztucznej inteligencji (system AI) oznacza oprogramowanie opracowane z wykorzystaniem jednej lub większej liczby technik i podejść wymienionych w załączniku I, które dla określonego zestawu celów zdefiniowanych przez człowieka może generować wyniki takie jak treści, prognozy, rekomendacje lub decyzje wpływające na

środowiska, z którymi wchodzi w interakcję”. Systemy AI można podzielić na dwie kategorie: modele klasyczne – tzw. słaba AI („top-down” AI), modele koneksjonistyczne – tzw. silna AI („bottom-up” AI). Silna AI obejmuje sieci modelowane na wzór ludzkiego mózgu. Słaba AI określana jest jako sztuczna inteligencja zadaniowa lub ilościowa. Natomiast silna AI wykazuje jakościową odmienność, przykładowo jest zdolna do wyboru własnych celów – w tym celów egzystencjalnych – oraz posiada swoistą formę świadomości. Do chwili obecnej nie opracowano systemu, który można byłoby uznać za silną AI. Wydaje się jednak, że jest to jedynie kwestia czasu. Naukowcy podkreślają, że pytanie dotyczące w pełni autonomicznej AI nie brzmi, czy stanie się ona elementem rzeczywistości, lecz kiedy to nastąpi. Autonomiczna AI niewątpliwie znajdzie zastosowanie w procesie pracy. Jej wykorzystanie zwiększy konkurencyjność oraz przyspieszy proces wytwarzania dóbr i świadczenia usług. Już obecnie testowane i wykorzystywane są maszyny półautonomiczne (np. samochody autonomiczne, drony, roboty medyczne, roboty do opieki nad osobami starszymi, roboty pracujące w kopalniach, roboty przeznaczone do dostarczania przesyłek).

Rewolucja technologiczna wyprzedza rozwiązania legislacyjne. Aby jednak zapobiec niebezpiecznej sytuacji, w której autonomiczna sztuczna inteligencja funkcjonowałaby poza obowiązującymi ramami prawnymi, już obecnie konieczne stało się rozstrzygnięcie problemu odpowiedzialności cywilnej za szkody wyrządzone w wyniku działania autonomicznej AI wykorzystywanej w procesie pracy. Mając na uwadze globalny zasięg sztucznej inteligencji oraz międzynarodowy charakter korporacji rozwijających technologie AI, rozwiązaniem optymalnym byłoby opracowanie i wdrożenie uniwersalnego modelu odpowiedzialności AI w skali światowej. Osiągnięcie tego celu jest jednak obecnie niemożliwe z kilku przyczyn. Po pierwsze, systemy prawne w poszczególnych częściach świata istotnie różnią się między sobą. Nie istnieją jednolite, ogólnoświatowe zasady odpowiedzialności cywilnej. Po drugie, mocarstwa technologiczne, takie jak Stany Zjednoczone oraz Chiny, prowadzą wyścig w zakresie opracowywania coraz bardziej zaawansowanych systemów sztucznej inteligencji oraz zdobywania globalnej hegemonii w tym obszarze. Państwa te konkurują na płaszczyźnie technologicznej i gospodarczej, co oznacza występowanie sprzecznych interesów. W konsekwencji rywalizacja w dziedzinie AI toczy się nie tylko pomiędzy poszczególnymi korporacjami, lecz także pomiędzy państwami i regionami świata. Z tego względu dążenie do wypracowania jednolitego modelu odpowiedzialności AI jest – przynajmniej na obecnym etapie – nierealne. Z drugiej strony zasadne i możliwe jest opracowanie regionalnego modelu odpowiedzialności AI na poziomie Unii Europejskiej. Realizacji tego celu poświęcony będzie niniejszy projekt. Dla jego osiągnięcia nie można jednak ograniczyć się wyłącznie do koncepcji i projektów aktów normatywnych odnoszących się do prawodawstwa Unii Europejskiej. Niezbędne jest również prześledzenie rozwiązań postulowanych w doktrynie w kontekście ustawodawstw Stanów Zjednoczonych, Chin oraz Singapuru jako państw wiodących w rozwoju technologii sztucznej inteligencji (Final Report 2021 of the National Security Commission on Artificial Intelligence). Pierwsze na świecie publiczne wdrożenie autonomicznych pojazdów miało miejsce w Singapurze.

Zgodnie z obowiązującym stanem prawnym pracodawca ponosi wobec osób trzecich odpowiedzialność za szkody wyrządzone przez pracownika, a więc przez człowieka. Z kolei autonomiczne systemy i maszyny oparte na sztucznej inteligencji były dotychczas kwalifikowane w prawie Unii Europejskiej jako produkty i obejmowane reżimem odpowiedzialności za produkt. Oznacza to, że odpowiedzialność za szkody wyrządzone przez AI przypisywana jest producentom towarów wyposażonych w systemy sztucznej inteligencji. Odpowiedzialność producenta za produkt ma charakter odpowiedzialności na zasadzie ryzyka. Ponadto sformułowano przesłanki odpowiedzialności operatorów systemów AI (Rezolucja Parlamentu Europejskiego z dnia 20 października 2020 r. zawierająca zalecenia dla Komisji w

sprawie systemu odpowiedzialności cywilnej za sztuczną inteligencję). Jako uzasadnienie odpowiedzialności operatora wskazano fakt, iż sprawuje on kontrolę nad ryzykiem związanym z funkcjonowaniem systemu AI. Podkreślono, że pojęcie operatora obejmuje zarówno operatora frontendowego, jak i backendowego, przy czym ten ostatni objęty jest zakresem dyrektywy w sprawie odpowiedzialności za produkt. Przez operatora frontendowego rozumie się każdą osobę fizyczną lub prawną, która sprawuje określony stopień kontroli nad ryzykiem związanym z działaniem i funkcjonowaniem systemu AI oraz czerpie korzyści z jego funkcjonowania. Natomiast operator backendowy oznacza każdą osobę fizyczną lub prawną, która w sposób ciągły określa cechy technologii oraz dostarcza dane i zasadnicze usługi wsparcia zaplecza technicznego, a tym samym również sprawuje określony stopień kontroli nad ryzykiem związanym z działaniem i funkcjonowaniem systemu AI. Zgodnie z projektem Rozporządzenia Parlamentu Europejskiego i Rady operator systemu AI wysokiego ryzyka ponosi odpowiedzialność na zasadzie ryzyka (odpowiedzialność ścisłą) za wszelką szkodę lub krzywdę wyrządzoną przez fizyczną lub wirtualną aktywność, urządzenie lub proces napędzany przez ten system AI. Jednocześnie operator systemu AI, który nie stanowi systemu wysokiego ryzyka, podlega odpowiedzialności opartej na zasadzie winy za wszelką szkodę lub krzywdę wyrządzoną przez fizyczną lub wirtualną aktywność, urządzenie lub proces napędzany przez dany system AI.

Obecnie stosowanie reżimu odpowiedzialności za produkt do sztucznej inteligencji nie jest uznawane za rozwiązanie optymalne. W przypadku szkód wyrządzonych przez AI nie jest łatwe jednoznaczne wskazanie producenta. W procesie tworzenia systemu sztucznej inteligencji zazwyczaj uczestniczy bowiem wielu różnych aktorów. Sytuacja dodatkowo się komplikuje, gdy możliwe jest skonstruowanie produktu końcowego z elementów pochodzących od różnych producentów, a sam proces ma charakter rozciągnięty w czasie. Należy również mieć na uwadze, że immanentną cechą autonomicznej AI jest jej zdolność do nieustannego uczenia się i w konsekwencji osiągania coraz wyższego poziomu zaawansowania. W rezultacie proces „wytwarzania” autonomicznej AI co do zasady nigdy nie zostaje definitywnie zakończony.

W świetle powyższego formułuję następujące pytania badawcze: jaka podstawa odpowiedzialności powinna znaleźć zastosowanie w sytuacji, gdy autonomiczna AI wyrządza szkodę osobie trzeciej podczas wykonywania czynności na rzecz określonego pracodawcy? Czy możliwe jest zastosowanie zasady winy, właściwej w odniesieniu do pracowników (ludzi), również wobec autonomicznej AI, a w szczególności czy w ogóle wykonalne jest ustalenie jakiegokolwiek „winy” autonomicznej sztucznej inteligencji? Rozważając alternatywną podstawę odpowiedzialności, tj. odpowiedzialność na zasadzie ryzyka, należy w szczególności zbadać, czy odpowiedzialność ta powinna – jak dotychczas – obciążać producenta albo operatora bądź pracodawcę, czy też zasadne byłoby przypisanie jej innemu podmiotowi. Kolejną kwestią wymagającą analizy jest to, czy pracodawca powinien zostać zwolniony z odpowiedzialności wobec osób trzecich za szkody wyrządzone przez autonomiczną AI, z której funkcjonowania czerpie korzyści, lecz która jest zdolna do podejmowania decyzji niezależnie od jego wiedzy lub zamiaru. Odpowiedź na to pytanie stanowić będzie punkt wyjścia do rozważań nad problematyką odpowiedzialności opartej na zasadzie słuszności. Udzielenie odpowiedzi na powyższe pytania badawcze pozwoli na sformułowanie wniosku, czy możliwe jest skuteczne rozwiązanie nowych, dotychczas nieznanych problemów związanych z AI przy wykorzystaniu jednego z istniejących modeli odpowiedzialności, a jeżeli tak – który z tych modeli jest najlepiej dostosowany do tego celu. Jeżeli natomiast analizowane reżimy odpowiedzialności okażą się nieadekwatne do zidentyfikowanych problemów związanych ze sztuczną inteligencją, w końcowym etapie badań przedstawiona zostanie propozycja nowego modelu odpowiedzialności dedykowanego autonomicznej AI.

2) Znaczenie projektu badawczego. W nauce prawa toczy się już dyskusja dotycząca modelu odpowiedzialności cywilnej za szkody wyrządzone przez działanie autonomicznych maszyn. W doktrynie przedstawiono różne propozycje uregulowania odpowiedzialności cywilnej za szkody spowodowane przez takie maszyny. Autorzy analizujący tę problematykę, po przeprowadzeniu stosownych badań, wyrazili pogląd, iż obowiązujące przepisy w badanych systemach prawnych są niewystarczające wobec nowego wyzwania technologicznego, jakim jest sztuczna inteligencja. Zgodnie z jedną z koncepcji (Jones, 2019), nowe ustawodawstwo w zakresie odpowiedzialności za autonomiczne maszyny mogłoby zostać oparte na znanym w wielu systemach prawnych reżimie odpowiedzialności za zwierzęta, które dana osoba hoduje lub którymi się posługuje. Na wzór odpowiedzialności za zwierzęta można byłoby wprowadzić regulację przewidującą, że osoba poszkodowana przez autonomiczną maszynę będzie uprawniona do jej zajęcia. Poszkodowanemu przysługiwałoby ustawowe prawo zastawu na zajętej autonomicznej maszynie w celu zabezpieczenia należnego odszkodowania. W razie niewypłacenia przez właściciela maszyny lub inny podmiot będący jej posiadaczem i czerpiący korzyści z jej działania należnego odszkodowania, poszkodowany mógłby sprzedać autonomiczną maszynę i z uzyskanej ceny zaspokoić przysługujące mu roszczenie. Zdaniem autora tej koncepcji model ten miałby uzasadnienie, ponieważ autonomiczne maszyny będą posiadały znaczną wartość rynkową, co w praktyce mogłoby umożliwić realne pokrycie szkody. Jednocześnie wskazano, że właściciele lub posiadacze autonomicznych maszyn mogliby zabezpieczać swoje interesy poprzez zawarcie odpowiedniej umowy ubezpieczenia odpowiedzialności cywilnej. Zgodnie z innym stanowiskiem (Villarica, 2017) odpowiedzialność za szkody wyrządzone przez autonomiczne maszyny powinna zostać uregulowana poprzez zastosowanie zasad odpowiedzialności za produkt. W konsekwencji, niezależnie od tego, kto jest właścicielem autonomicznej maszyny i kto korzysta z jej działania, odpowiedzialność za szkody powinna obciążać podmiot, który maszynę wyprodukował i wprowadził do obrotu. W tej koncepcji przyjmuje się, że wypadek i wynikająca z niego szkoda są zawsze następstwem błędu projektowego, wady algorytmu stanowiącego podstawę funkcjonowania maszyny lub jej konstrukcyjnej wady. W rezultacie odpowiedzialność cywilna nie powinna obciążać żadnego innego podmiotu niż producent autonomicznej maszyny albo podmiot wprowadzający ją do obrotu – jeżeli jest nim podmiot inny niż producent. Według jeszcze innej koncepcji uwaga ustawodawcy powinna koncentrować się wyłącznie na twórcach algorytmów, na podstawie których autonomiczna maszyna jest zdolna do podejmowania samodzielnych decyzji. Zgodnie z tym poglądem odpowiedzialność powinni ponosić wyłącznie programiści tworzący odpowiednie algorytmy. Jeżeli określony patent technologiczny zostaje sprzedany producentowi, ani producent, ani tym bardziej przyszły nabywca autonomicznej maszyny (właściciel) nie są w stanie przewidzieć decyzji, jakie maszyna podejmie w oparciu o algorytmy znane jedynie ich twórcom. W konsekwencji producenci oraz przyszli właściciele autonomicznych maszyn nie powinni ponosić odpowiedzialności za ich działanie, gdyż działanie to może być w pełni nieprzewidywalne. Zwrócono jednak uwagę, że wprowadzenie wyłącznej – lub nawet częściowej – odpowiedzialności programistów mogłoby prowadzić do tzw. „efektu mrożącego”, polegającego na spowolnieniu procesu tworzenia przełomowych rozwiązań technologicznych z obawy przed odpowiedzialnością za skutki ich praktycznego zastosowania. W konsekwencji efekt ten mógłby negatywnie oddziaływać na rozwój i wdrażanie innowacyjnych technologii. Ostatnia z prezentowanych koncepcji zakłada, że rewolucyjne zmiany technologiczne wymagają równie odważnych reform obowiązujących ram prawnych. Wskazuje się, iż moment wprowadzenia do obrotu w pełni autonomicznych maszyn będzie właściwym czasem na przyznanie im podmiotowości prawnej oraz ukształtowanie koncepcji tzw. cyfrowej osoby prawnej. Zgodnie z tym ujęciem autonomiczne maszyny o określonym poziomie zaawansowania technologicznego, działające na podstawie własnych, niezależnych decyzji i niepodążające za decyzjami innych osób, byłyby wpisywane do

specjalnego rejestru cyfrowych osób prawnych. Z chwilą wpisu do rejestru uzyskiwałyby osobowość prawną i stawałyby się cyfrowymi osobami prawnymi. Wiązałoby się to z obowiązkiem wykonywania nałożonych na nie obowiązków, w tym w szczególności z obowiązkiem zawarcia umowy ubezpieczenia odpowiedzialności cywilnej. W rezultacie właściciele (współwłaściciele, udziałowcy) cyfrowej osoby prawnej byłiby zobowiązani do ponoszenia kosztów ubezpieczenia autonomicznej maszyny w formie składki, lecz nie ponosiliby bezpośredniej odpowiedzialności za ewentualne szkody przez nią wyrządzone. Odpowiedzialność spoczywałaby wyłącznie na cyfrowej osobie prawnej. Analogiczne rozwiązanie funkcjonuje obecnie w odniesieniu do osób prawnych – odpowiedzialność ponosi np. spółka z ograniczoną odpowiedzialnością, a nie jej wspólnicy będący osobami fizycznymi. Poszczególne koncepcje dotyczące ukształtowania odpowiedzialności cywilnej AI różnią się między sobą m.in. tym, że niektóre z nich mają charakter zachowawczy (np. traktowanie AI w ramach odpowiedzialności za produkt), inne natomiast są wysoce innowacyjne (np. samodzielna odpowiedzialność AI jako cyfrowej osoby prawnej).

Potrzebę wypracowania adekwatnych reżimów prawnych w odniesieniu do autonomicznych maszyn dostrzegło wiele państw, w szczególności te, które aktywnie uczestniczą w globalnym wyścigu technologicznym. Problematyka ta stanowi również przedmiot zainteresowania Unii Europejskiej. W dokumencie Komisji Europejskiej z dnia 19 lutego 2020 r. pt. „Biała księga w sprawie sztucznej inteligencji – europejskie podejście do doskonałości i zaufania” przypomniano, że zgodnie z prawem UE odpowiedzialność za produkt wadliwy ponosi producent, co dotyczy także produktów opartych na AI. Szerokie rozważania dotyczące bezpieczeństwa i odpowiedzialności w kontekście AI, Internetu Rzeczy (IoT) oraz robotyki zawarto w Sprawozdaniu Komisji z dnia 19 lutego 2020 r. pt. „Sprawozdanie w sprawie skutków dla bezpieczeństwa i odpowiedzialności związanych ze sztuczną inteligencją, Internetem Rzeczy i robotyką”. W Sprawozdaniu wskazano, że jedną z podstawowych cech produktów i systemów opartych na AI jest ich nietransparentność, wynikająca m.in. ze zdolności do ciągłego doskonalenia poprzez uczenie się na podstawie zdobytych doświadczeń. W zależności od przyjętej metodologii produkty i systemy AI mogą cechować się różnym stopniem nietransparentności. Wykorzystywanie ogromnych zbiorów danych, operowanie na algorytmach oraz brak przejrzystości procesu decyzyjnego podważają przewidywalność działania systemów AI oraz możliwość ustalenia przyczyn powstania szkody. W tym kontekście proces decyzyjny autonomicznego systemu opartego na AI może być trudny do prześledzenia (tzw. efekt „czarnej skrzynki”). W przyszłości może dojść do eskalacji sytuacji, w których nie będzie możliwe uprzednie określenie wszystkich rezultatów działania AI. Proces samodzielnego uczenia się systemów AI może prowadzić do podejmowania decyzji odbiegających od pierwotnych założeń producenta oraz oczekiwań użytkowników. Podkreślono, że choć człowiek nie musi rozumieć każdego etapu procesu decyzyjnego AI, konieczne jest, aby możliwe było wyjaśnienie, w jaki sposób – przy użyciu algorytmu – system podjął daną decyzję. Ma to szczególne znaczenie w kontekście ustalania odpowiedzialności. Zauważono również, że obowiązujące w UE przepisy dotyczące bezpieczeństwa produktów nie odnoszą się w sposób wystarczająco jasny do rosnących zagrożeń wynikających z nietransparentności systemów AI. W związku z tym zasadne jest rozważenie wprowadzenia szczególnych wymogów w zakresie przejrzystości algorytmów oraz nadzoru człowieka nad ich funkcjonowaniem w całym cyklu życia produktu. Dodatkowo proponuje się nałożenie na twórców algorytmów obowiązku ujawnienia parametrów projektowych oraz metadanych w przypadku szkody wyrządzonej przez AI.

Komisja Europejska zwróciła także uwagę na złożoność środowiska Internetu Rzeczy (IoT), w którym współdziałają liczne urządzenia i usługi oparte na AI. Z uwagi na tę złożoność poszkodowanym może być bardzo trudno ustalić podmiot odpowiedzialny oraz spełnić wymogi

dowodowe przewidziane w prawie krajowym. Autonomia i nietransparentność AI dodatkowo utrudniają ustalenie, czy szkoda wynikała z działania człowieka, co miałoby znaczenie dla odpowiedzialności opartej na winie. Wskazano ponadto na potrzebę dostosowania krajowych przepisów dotyczących ciężaru dowodu w sprawach, w których szkoda została wyrządzona przez produkt lub system oparty na AI. Zaproponowano powiązanie ciężaru dowodu z przestrzeganiem przez operatorów określonych obowiązków, np. w zakresie cyberbezpieczeństwa, oraz jego przerzucenie na operatora w razie ich naruszenia.

Autonomia i nietransparentność AI, IoT i robotyki podważają skuteczność części unijnych i krajowych regulacji odpowiedzialności. Jednocześnie niezwykle istotne jest zapewnienie osobom poszkodowanym przez systemy AI ochrony prawnej na poziomie nie niższym niż w przypadku szkód wyrządzonych bez udziału AI. Równie ważne jest, aby przedsiębiorcy rozwijający nowe technologie znali zasady swojej potencjalnej odpowiedzialności cywilnej oraz mogli odpowiednio zabezpieczyć się przed ryzykiem, np. poprzez zawarcie stosownej umowy ubezpieczenia.

Komisja zwróciła również uwagę, że najnowsze technologie oparte na AI mogą w określonych sytuacjach stanowić zagrożenie dla zdrowia psychicznego człowieka, np. w warunkach współpracy z humanoidalnymi robotami. Okoliczność ta powinna zostać uwzględniona przy definiowaniu pojęcia bezpieczeństwa produktu. Pojawiła się także propozycja zobowiązania producentów robotów humanoidalnych do uwzględniania szkód niemajątkowych wyrządzanych przez autonomiczne maszyny.

Podkreślono, że zgodnie z zasadami prawa UE to producent odpowiada za zapewnienie bezpieczeństwa produktu w całym jego cyklu życia. Jednocześnie zauważono, że systemy AI mogą być modyfikowane już po wprowadzeniu do obrotu. W związku z tym wskazano na potrzebę redefinicji pojęcia „wprowadzenia do obrotu” w dyrektywie 85/374/EWG w sprawie odpowiedzialności za produkty wadliwe, z uwzględnieniem specyfiki systemów autonomicznych. Autonomia może bowiem istotnie zmieniać właściwości produktu, w tym jego bezpieczeństwo. Nie jest zatem jasne, czy w dłuższej perspektywie producent powinien ponosić odpowiedzialność za skutki samouczenia się w pełni autonomicznych maszyn, ani w jakim zakresie jest w stanie przewidzieć zachodzące w nich zmiany. Zdaniem Komisji Europejskiej niektóre autonomiczne urządzenia i systemy oparte na AI mogą charakteryzować się szczególnym profilem ryzyka odpowiedzialności, zwłaszcza gdy mogą powodować szkody na osobie lub mieniu, a nawet narażać ogół społeczeństwa na niebezpieczeństwo (np. w przypadku autonomicznych pojazdów, dronów czy systemów zarządzania ruchem). W ocenie Komisji problemy wynikające z autonomii i nietransparentności AI mogą zostać rozwiązane poprzez zastosowanie odpowiedzialności na zasadzie ryzyka. Dzięki takiemu modelowi poszkodowany mógłby dochodzić odszkodowania niezależnie od wykazania winy.

Z drugiej strony Parlament Europejski w rezolucji z dnia 16 lutego 2017 r. zawierającej zalecenia dla Komisji w sprawie przepisów prawa cywilnego dotyczących robotyki (2015/2103(INL)) podkreślił, że przyszły instrument prawny w zakresie odpowiedzialności powinien zostać oparty na szczegółowej ocenie przeprowadzonej przez Komisję, mającej na celu ustalenie, czy właściwe będzie zastosowanie odpowiedzialności na zasadzie ryzyka (odpowiedzialności ścisłej), czy też podejścia opartego na zarządzaniu ryzykiem. Parlament zwrócił uwagę na konieczność rozważenia, czy w odniesieniu do odpowiedzialności za działanie autonomicznych robotów nie zachodzi potrzeba wprowadzenia nowych zasad i norm prawnych. Odpowiedź na to pytanie ma kluczowe znaczenie, w szczególności w sytuacjach, w których nie jest możliwe prześledzenie przyczyn szkody, a odpowiedzialność nie może zostać przypisana konkretnej osobie fizycznej. W rezolucji wskazano, że obowiązujące reżimy odpowiedzialności mogą okazać się niewystarczające w odniesieniu do szkód wyrządzonych

przez roboty nowej generacji, wyposażone w pełną autonomię, zdolność adaptacji do środowiska oraz uczenia się. Cechy te wiążą się z określoną nieprzewidywalnością ich zachowań. Roboty uczą się na podstawie własnych, zróżnicowanych doświadczeń i wchodzą w interakcje ze środowiskiem w sposób unikalny i trudny do przewidzenia. Wskazuje to na potrzebę wprowadzenia wspólnych definicji unijnych dotyczących systemów cyberfizycznych, systemów autonomicznych, inteligentnych robotów autonomicznych oraz ich podkategorii, z uwzględnieniem takich cech jak: uzyskiwanie autonomii dzięki czujnikom lub wymianie danych ze środowiskiem; zdolność do samouczenia się; posiadanie przynajmniej minimalnej formy fizycznej; dostosowywanie zachowania do środowiska; brak funkcji życiowych w sensie biologicznym. Parlament Europejski wskazał także na potrzebę stworzenia obowiązkowego systemu ubezpieczeniowego, w ramach którego producenci autonomicznych robotów byłoby zobowiązani do zawierania umów ubezpieczenia odpowiedzialności cywilnej. System ten mógłby zostać uzupełniony specjalnym funduszem służącym kompensacji szkód nieobjętych ochroną ubezpieczeniową. Zaproponowano również utworzenie kompleksowego unijnego systemu rejestracji zaawansowanych robotów autonomicznych. Każdemu robotowi przypisany byłby indywidualny numer rejestracyjny powiązany z funduszem kompensacyjnym, co umożliwiłoby uzyskanie informacji o ewentualnych ograniczeniach odpowiedzialności oraz o funduszu uzupełniającym system ubezpieczenia. W dłuższej perspektywie rozważano możliwość przyznania w pełni autonomicznym robotom statusu „osób elektronicznych” odpowiedzialnych za wyrządzone szkody.

Zgodnie z rezolucją Parlamentu Europejskiego z dnia 20 października 2020 r. w sprawie systemu odpowiedzialności cywilnej za sztuczną inteligencję (2020/2014(INL)), rozwój AI stanowi istotne wyzwanie dla istniejących ram odpowiedzialności. Powszechne wykorzystywanie systemów AI w życiu codziennym może prowadzić do sytuacji, w których ich nietransparentność (efekt „czarnej skrzynki”) oraz wielość podmiotów uczestniczących w ich cyklu życia uczynią identyfikację podmiotu kontrolującego ryzyko niezwykle kosztowną lub wręcz niemożliwą. Trudności te potęguje łączność systemów AI z innymi systemami, zależność od danych zewnętrznych, podatność na zagrożenia cyberbezpieczeństwa oraz rosnąca autonomia wynikająca z zastosowania uczenia maszynowego i głębokiego uczenia. Systemy AI mogą być również wykorzystywane w sposób prowadzący do poważnych naruszeń godności ludzkiej i wartości europejskich, np. poprzez nieuprawnione śledzenie osób, wprowadzanie systemów oceny społecznej, podejmowanie stronniczych decyzji w sprawach ubezpieczeń zdrowotnych, kredytów czy zatrudnienia, a także konstruowanie autonomicznych systemów uzbrojenia. W takich przypadkach ustalenie związku między szkodą a konkretną decyzją człowieka może być szczególnie utrudnione. Operator może twierdzić, że szkoda wynikała z autonomicznego działania systemu pozostającego poza jego kontrolą. Jednocześnie samo funkcjonowanie autonomicznego systemu AI nie może automatycznie stanowić podstawy odpowiedzialności. Może to prowadzić do sytuacji, w których rozkład odpowiedzialności będzie niesprawiedliwy lub nieefektywny, a poszkodowany pozostanie bez kompensacji.

Rezolucja podkreśla jednak, że podmiot, który tworzy, utrzymuje, kontroluje lub ingeruje w system AI, powinien ponosić odpowiedzialność za szkody przez niego wyrządzone. Wynika to z ogólnych zasad sprawiedliwości, zgodnie z którymi podmiot stwarzający ryzyko dla społeczeństwa powinien je minimalizować, a w razie jego materializacji – naprawić powstałą szkodę. Zdaniem Parlamentu rozwój AI nie wymaga całkowitej reformy systemów odpowiedzialności, lecz raczej ich dostosowania poprzez wprowadzenie celowanych zmian i nowych przepisów, przy jednoczesnym zachowaniu dotychczasowych, sprawnie funkcjonujących reżimów.

Komisja Europejska przedstawiła wniosek dotyczący Rozporządzenia Parlamentu Europejskiego i Rady ustanawiającego zharmonizowane przepisy dotyczące sztucznej

inteligencji (Artificial Intelligence Act) oraz zmieniającego niektóre akty ustawodawcze Unii (Bruksela, 21.4.2021). Projektując nowe regulacje, Komisja przyjęła podejście oparte na analizie ryzyka. Zakres reżimu prawnego dotyczącego systemów sztucznej inteligencji opiera się na zasadzie: „im większe ryzyko związane z AI, tym bardziej rygorystyczny reżim prawny”. W konsekwencji przedmiotem regulacji nie są systemy sztucznej inteligencji jako takie, lecz sposób ich wykorzystania. W rezultacie systemy AI zostały podzielone na cztery kategorie, według poziomu ryzyka związanego z ich konkretnym zastosowaniem (ryzyko minimalne, ograniczone, wysokie, niedopuszczalne).

Ryzyko minimalne lub zerowe – systemy AI, których użycie nie wiąże się z żadnym ryzykiem albo ryzyko to jest znikome – nie będą objęte projektowanym reżimem prawnym i będą mogły być rozwijane oraz wykorzystywane bez ograniczeń. Do tej kategorii zalicza się większość obecnie stosowanych rozwiązań z zakresu sztucznej inteligencji, takich jak filtry antyspamowe, systemy przekształcania mowy na tekst czy automatyczne tłumaczenia tekstów.

Ryzyko ograniczone – systemy AI przeznaczone do interakcji z człowiekiem – będą podlegały wyłącznie obowiązkom informacyjnym oraz wymogom przejrzystości, tak aby użytkownik miał świadomość, że wchodzi w interakcję z systemem AI, a nie z człowiekiem. Przykładem takich systemów są chatboty.

Ryzyko wysokie – obejmuje systemy mogące w istotny sposób ingerować w życie człowieka. Systemy te będą podlegały szeregowi surowych obowiązków adresowanych do ich producentów, dystrybutorów oraz użytkowników. Do systemów wysokiego ryzyka zaliczyć można w szczególności systemy zdalnej identyfikacji biometrycznej w czasie rzeczywistym, systemy wykorzystywane w procesach rekrutacyjnych lub przy ocenie zdolności kredytowej, a także zastosowania medyczne AI.

Ryzyko niedopuszczalne – wykorzystywanie systemów AI zakwalifikowanych jako „niedopuszczalne”, ze względu na ich sprzeczność z prawami podstawowymi Unii Europejskiej, będzie zakazane. Chodzi tu przykładowo o systemy służące do oceny wiarygodności społecznej (social scoring) lub rozwiązania wykorzystujące manipulację podprogową w celu wpływania na zachowania ludzi.

W projekcie Rozporządzenia najwięcej miejsca poświęcono systemom AI wysokiego ryzyka. Producenci takich systemów będą zobowiązani do realizacji szeregu obowiązków mających na celu minimalizację ryzyka związanego z danymi rozwiązaniami. W szczególności będą zobowiązani do:

- wykorzystywania danych wysokiej jakości do trenowania i testowania systemów AI,
- sporządzania dokumentacji technicznej opisującej system AI oraz zawierającej informacje niezbędne do oceny zgodności systemu z wymogami Rozporządzenia,
- zapewnienia prowadzenia rejestrów w trakcie funkcjonowania systemów AI,
- zapewnienia odpowiedniego poziomu przejrzystości systemu (tj. umożliwienia użytkownikom interpretowania skutków działania systemu oraz właściwego korzystania z tych skutków),
- opracowania instrukcji dla użytkowników, określającej zarówno możliwości, jak i ograniczenia systemu AI,
- zapewnienia możliwości nadzoru człowieka nad funkcjonowaniem systemu AI,
- zapewnienia odpowiedniego poziomu odporności na błędy,

- zapewnienia cyberbezpieczeństwa oraz odporności na awarie lub ataki hakerskie.

Również importerzy, dystrybutorzy oraz użytkownicy systemów AI wysokiego ryzyka będą obciążeni szeregiem obowiązków – od weryfikacji prawidłowości wprowadzenia systemu do obrotu po monitorowanie jego funkcjonowania.

Projekt Rozporządzenia przewiduje powołanie nowego organu nadzorczego – Europejskiej Rady ds. Sztucznej Inteligencji. W jej skład wchodzić będą kierownicy krajowych organów nadzorczych właściwych w sprawach sztucznej inteligencji z poszczególnych państw członkowskich oraz Europejski Inspektor Ochrony Danych. Do zadań Rady należeć będzie między innymi konsolidowanie wiedzy eksperckiej oraz jej upowszechnianie wśród państw członkowskich, a także wydawanie opinii i zaleceń w sprawach związanych z wdrażaniem Rozporządzenia.

Uzasadnieniem podjęcia przedstawionego problemu badawczego jest fakt, iż ustawodawca unijny nie zdecydował się na przyznanie sztucznej inteligencji osobowości prawnej w połączeniu z jej samodzielną odpowiedzialnością cywilną za wyrządzone szkody. Obecnie AI klasyfikowana jest nie jako cyfrowa osoba prawna, lecz jako produkt, a ustawodawca skoncentrował się na reżimie odpowiedzialności za produkt oraz na odpowiedzialności operatorów systemów AI. Powstaje zatem pytanie, czy ustawodawstwo Unii Europejskiej nie ma charakteru nadmiernie zachowawczego oraz czy umożliwi ono budowanie pełnego zaufania do AI i jej szerokie wykorzystanie w procesie pracy. Unia Europejska aspiruje do roli lidera w zakresie rozwoju w pełni autonomicznej sztucznej inteligencji. Cel ten zamierza osiągnąć między innymi poprzez tworzenie regulacji prawnych kompatybilnych z rozwojem nowych technologii AI. W tym kontekście konieczne jest zbadanie proponowanych rozwiązań w zakresie odpowiedzialności za AI przyjmowanych w ustawodawstwach światowych liderów w dziedzinie technologii AI, takich jak Stany Zjednoczone, Chiny czy Singapur. Propozycje formułowane w literaturze naukowej odnoszącej się do zagranicznych systemów prawnych mogą być pomocne w ocenie, czy kierunek obrany przez Unię Europejską (odpowiedzialność za AI = odpowiedzialność za produkt) jest właściwy. Jeżeli tak nie jest, mogą one wskazać możliwe alternatywne rozwiązania. Wyniki przeprowadzonych badań pozwolą udzielić odpowiedzi na pytanie, w jaki sposób najnowsze technologie zmieniają proces pracy oraz jak – w tym kontekście – należy ukształtować odpowiedzialność pracodawcy, aby wykorzystanie nowoczesnych technologii było możliwe i następowało z korzyścią dla społeczeństwa. Badania te mogą stanowić punkt odniesienia dla dalszych analiz oraz przyczynią się do rozwoju dyskusji naukowej. Niezbędna jest pogłębiona debata naukowa w celu udzielenia odpowiedzi na pytanie, jaki model legislacyjny należy wprowadzić, aby zapewnić pełne i optymalne wykorzystanie innowacyjnych technologii, przy jednoczesnym zagwarantowaniu ochrony praw człowieka. Projekt badawczy różni się istotnie od dotychczasowych badań nad cywilnoprawną odpowiedzialnością autonomicznej AI, gdyż nowa perspektywa badawcza polega na zestawieniu obowiązujących i projektowanych regulacji unijnych z propozycjami doktryny prawa odnoszącymi się do systemów prawnych najbardziej zaawansowanych technologicznie państw, takich jak USA, Chiny czy Singapur.

3) Koncepcja i plan badań. Badania zostaną przeprowadzone według następującego planu:

1. Przegląd i analiza propozycji dotyczących ukształtowania odpowiedzialności AI w ustawodawstwie Stanów Zjednoczonych.
2. Przegląd i analiza koncepcji dotyczących odpowiedzialności AI w ustawodawstwie chińskim.

3. Przegląd i analiza propozycji dotyczących odpowiedzialności AI w ustawodawstwie Singapuru.
4. Przegląd i analiza proponowanych w literaturze naukowej rozwiązań dotyczących odpowiedzialności AI w prawie Unii Europejskiej.
5. Przegląd i analiza obowiązujących oraz projektowanych instrumentów prawnych UE w zakresie odpowiedzialności za AI.
 - 5.1. Odpowiedź na pytanie: jakie korzyści wynikają z przyjęcia takiej regulacji?
 - 5.2. Odpowiedź na pytanie: jakie komplikacje wiążą się z taką regulacją?
 - 5.3. Odpowiedź na pytanie, czy analizowane reżimy prawne są wystarczające i adekwatne wobec dynamicznego rozwoju technologii AI.
6. Zestawienie wniosków wynikających z badań przeprowadzonych w punktach 1–3 z wnioskami z badań wskazanych w punkcie 4.
7. Odpowiedź na pytanie, czy możliwe jest optymalne ukształtowanie odpowiedzialności za AI przy wykorzystaniu jednego z istniejących modeli odpowiedzialności, a jeżeli tak – który z nich jest najbardziej adekwatny.
8. W przypadku stwierdzenia nieadekwatności istniejących modeli odpowiedzialności do zidentyfikowanych problemów związanych z AI – sformułowanie propozycji nowego modelu odpowiedzialności dedykowanego autonomicznej AI.

Wyniki badań zostaną opublikowane w czasopiśmie naukowych o zasięgu międzynarodowym oraz zaprezentowane podczas międzynarodowych konferencji. Planowane badania mają charakter wyłącznie teoretyczny, co oznacza, iż ich realizacja nie wiąże się ze szczególnym ryzykiem. Pandemia wirusa SARS-CoV-2 spowolniła prace legislacyjne dotyczące optymalnego ukształtowania odpowiedzialności za AI na poziomie Unii Europejskiej. Stan ten ma jednak charakter przejściowy i należy oczekiwać, że w krótszej lub dłuższej perspektywie działania w tym obszarze ulegną intensyfikacji. Okoliczność ta nie stanowi zagrożenia dla planowanych badań, gdyż na bieżąco monitoruję wszystkie projekty legislacyjne i działania podejmowane na poziomie UE. Jednocześnie sformułowany temat badawczy jest niezwykle aktualny i dotyczy obszaru AI oraz technologii autonomicznych, który rozwija się dynamicznie. Oznacza to, że w najbliższym czasie można spodziewać się istotnego wzrostu liczby projektów badawczych oraz publikacji naukowych poświęconych tej problematyce. Ryzyko to może zostać zminimalizowane poprzez wypracowanie efektywnej metody selekcji publikacji oraz uwzględnianie w analizach jedynie wybranych, najbardziej relevantnych opracowań.

4) Metodologia badań. W pierwszej kolejności dokonam przeglądu koncepcji dotyczących ukształtowania odpowiedzialności AI w ustawodawstwie Stanów Zjednoczonych. Wybrane propozycje przedstawione w literaturze naukowej zostaną przeanalizowane między innymi pod kątem ich skuteczności w rozwiązywaniu istniejących problemów w zakresie odpowiedzialności za AI rozumianej jako odpowiedzialność za produkt. Analizowane koncepcje nie muszą stanowić gotowych rozwiązań dla prawa unijnego, lecz będą stanowić punkt odniesienia dla dalszych badań. Analogiczny proces badawczy zostanie przeprowadzony w odniesieniu do rozwiązań proponowanych w zakresie odpowiedzialności za AI w systemach prawnych Chin i Singapuru. Badania obejmą wyłącznie idee formułowane w kontekście rozwiązywania problemu odpowiedzialności za AI. Nie będą natomiast analizowane konkretne akty normatywne, ponieważ systemy prawne USA, Chin i Singapuru różnią się zasadniczo od systemu prawnego Unii Europejskiej i nie nadają się do bezpośredniej adaptacji. Sformułowanie ogólnych, a następnie bardziej szczegółowych wniosków dotyczących

cywilnoprawnej odpowiedzialności za AI pozwoli ustalić najnowsze koncepcje w tym zakresie w systemach prawnych USA, Chin i Singapuru. Ustalenia te będą kluczowe dla kolejnego etapu badań, skoncentrowanego na Unii Europejskiej. W pierwszej kolejności dokonam przeglądu propozycji dotyczących ukształtowania odpowiedzialności za AI w prawie UE, prezentowanych w literaturze naukowej. Następnie przeprowadzę analizę obowiązujących oraz projektowanych instrumentów prawnych UE w zakresie odpowiedzialności za AI. W przeciwieństwie do badań dotyczących USA, Chin i Singapuru, w odniesieniu do Unii Europejskiej analizie poddane zostaną również akty prawne, co wynika z celu projektu, jakim jest weryfikacja i zaproponowanie modelu odpowiedzialności za AI możliwego do zastosowania w prawie UE. Na kolejnym etapie wnioski z badań przeprowadzonych w punktach 1–3 zostaną zestawione z wnioskami wynikającymi z badań dotyczących UE. Takie zestawienie umożliwi udzielenie odpowiedzi na pytanie, czy możliwe jest optymalne ukształtowanie odpowiedzialności za AI przy wykorzystaniu jednego z istniejących modeli odpowiedzialności, a jeżeli tak – który z nich jest najbardziej adekwatny. Jeżeli jednak wyniki wcześniejszych etapów badań wykażą, że istniejące modele odpowiedzialności są nieadekwatne wobec zidentyfikowanych problemów związanych z AI, w końcowej fazie prac badawczych zostanie sformułowana propozycja nowego modelu odpowiedzialności dedykowanego autonomicznej sztucznej inteligencji.