



KRAJOWA IZBA
RADCÓW PRAWNYCH

Dokument dotyczy dwóch
projektów rozporządzeń – Digital
Omnibus i Digital Omnibus on AI

Digital Omnibus

Analiza wybranych
zagadnień

Warszawa, luty 2026

Analiza wybranych aspektów projektu pakietu zmian legislacyjnych Digital Omnibus oraz Digital Omnibus on AI



KRAJOWA IZBA
RADCÓW PRAWNYCH

Warszawa, 10 lutego 2026

Spis treści

I. Przedmiot i zakres analizy	3
II. Słowniczek	5
III. Uwagi ogólne	7
IV. Uwagi Szczegółowe	9
1. Zmiany w RODO (Rozporządzenie 2016/679)	9
1.1. Uwagi szczegółowe	9
1.2. Uwagi szczegółowe	11
2. Zmiany w AIA (Rozporządzenie 2024/1689)	21
2.1. Uwagi ogólne	21
2.2. Uwagi szczegółowe	22
V. Podsumowanie i wnioski	33

I. Przedmiot i zakres analizy

Celem niniejszego opracowania jest przedstawienie wybranych zastrzeżeń, uwag i propozycji kierunków zmian do wybranych projektów rozporządzeń zmieniających wchodzących w skład pakietu Digital Omnibus i Digital Omnibus on AI przedstawionych przez Komisję Europejską 19 listopada 2025 r.

Opracowanie sporządzone zostało przez następujących ekspertów Komisji Nowych Technologii Krajowej Rady Radców Prawnych:

- r. pr. dr hab. Dominik Lubasz
- r. pr. dr Michał Araszkiewicz
- r. pr. Anna Augustyn

Nie obejmuje ono kompleksowej analizy projektów rozporządzeń wchodzących w skład pakietu Digital Omnibus i Digital Omnibus on AI ani nie wskazuje na wszystkie zidentyfikowane problemy, wątpliwości interpretacyjne czy możliwe rozwiązania itp. Uwagi, wnioski i propozycje zawarte w opracowaniu są wynikiem dyskusji i rozważań powyżej wskazanych ekspertów.

Na wstępie autorzy chcą podkreślić, że popierają niezbędne działania legislacyjne na rzecz rozwoju innowacji technologicznej, w tym rozwoju sztucznej inteligencji i dostrzegają w tym aspekcie warunek konieczny do zachowania konkurencyjności gospodarczej Europy. Jednocześnie innowacyjność nie powinna być rozwijana kosztem ochrony praw jednostek, pewności prawa czy rozliczalności podmiotów działających na rynku AI, a korzyści płynące ze zmian przepisów powinny przewyższać potencjalne ryzyka w postaci osłabienia mechanizmów ochronnych. W ocenie autorów niektóre z proponowanych przez pakiet Digital Omnibus zmian budzą uzasadnione wątpliwości co do zachowania tej równowagi. Ponadto, niektóre spośród projektowanych rozwiązań budzą istotne zastrzeżenia z punktu widzenia zasad techniki prawodawczej.

Uwzględniając rolę samorządu radców prawnych jako zawodu zaufania publicznego w ochronie praw jednostek oraz interesu publicznego, konieczne jest zwrócenie uwagi na praktyczne skutki projektowanych zmian. Radcowie prawni mają bezpośredni kontakt z regulacjami dotyczącymi AI: zarówno w ramach własnej praktyki zawodowej (w tym przy

wdrażaniu i wykorzystywaniu narzędzi AI w pracy kancelarii), jak i przy świadczeniu usług prawnych na rzecz podmiotów rozwijających oraz stosujących systemy AI, w tym we wspieraniu organizacji w zapewnianiu zgodności z prawem, a także występują jako pełnomocnicy osób, których prawa mogą zostać naruszone w wyniku działania tych systemów. Z tego względu kluczowe znaczenie mają pewność prawa (w tym przewidywalność harmonogramu stosowania obowiązków AIA), przejrzystość regulacji oraz utrzymanie skutecznych mechanizmów ochrony praw podstawowych, warunkujących m.in. realną ochronę prawną osób fizycznych i prawidłowe funkcjonowanie wymiaru sprawiedliwości.

II. Słowniczek

Źródła Prawa	
AIA, AI Act lub Rozporządzenie 2024/1689	ROZPORZĄDZENIE PARLAMENTU EUROPEJSKIEGO I RADY (UE) 2024/1689 z dnia 13 czerwca 2024 r. w sprawie ustanowienia zharmonizowanych przepisów dotyczących sztucznej inteligencji oraz zmiany rozporządzeń (WE) nr 300/2008, (UE) nr 167/2013, (UE) nr 168/2013, (UE) 2018/858, (UE) 2018/1139 i (UE) 2019/2144 oraz dyrektyw 2014/90/UE, (UE) 2016/797 i (UE) 2020/1828 (akt w sprawie sztucznej inteligencji)
RODO lub Rozporządzenie 2016/679	ROZPORZĄDZENIE PARLAMENTU EUROPEJSKIEGO I RADY (UE) 2016/679 z dnia 27 kwietnia 2016 r. w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych i w sprawie swobodnego przepływu takich danych oraz uchylenia dyrektywy 95/46/WE (ogólne rozporządzenie o ochronie danych) (Dz. Urz. UE L 119 z 4.05.2016 r.)
Digital Omnibus lub Akt zbiorczy prawa cyfrowego	Wniosek ROZPORZĄDZENIE PARLAMENTU EUROPEJSKIEGO I RADY w sprawie zmiany rozporządzeń (UE) 2016/679, (UE) 2018/1724, (UE) 2018/1725, (UE) 2023/2854 oraz dyrektyw 2002/58/WE, (UE) 2022/2555 i (UE) 2022/2557 w odniesieniu do uproszczenia cyfrowych ram prawnych oraz uchylenia rozporządzeń (UE) 2018/1807, (UE) 2019/1150, (UE) 2022/868 i dyrektywy (UE) 2019/1024 (akt zbiorczy prawa cyfrowego), COM(2025) 837 final
Digital Omnibus on AI lub Akt zbiorczy prawa cyfrowego dotyczący AI	Wniosek ROZPORZĄDZENIE PARLAMENTU EUROPEJSKIEGO I RADY zmieniające rozporządzenia (UE) 2024/1689 i (UE) 2018/1139 w odniesieniu do uproszczenia wdrażania zharmonizowanych

	przepisów dotyczących sztucznej inteligencji (akt zbiorczy prawa cyfrowego dotyczący AI), COM(2025) 836 final
TFUE	Traktat o Funkcjonowaniu Unii Europejskiej (wersja skonsolidowana Dz. Urz. UE L 326)
Inne skróty	
FRIA	Fundamental Rights Impact Assessment (Ocena skutków systemów AI wysokiego ryzyka dla praw podstawowych)
GPAI	General Purpose AI (Sztuczna Inteligencja Ogólnego Przeznaczenia)
LIA	Legitimate Interest Assessment (test uzasadnionego interesu)
SSI lub system SI	System Sztucznej Inteligencji
ETPC	Europejski Trybunał Praw Człowieka
EDPB lub EROD	Europejska Rada Ochrony Danych (European Data Protection Board)
EDPS lub EIOD	Europejski Inspektor Ochrony Danych (European Data Protection Supervisor)
TS/TSUE	Trybunał Sprawiedliwości/ Trybunał Sprawiedliwości Unii Europejskiej (od dnia wejścia w życie Traktatu z Lizbony, to jest 1 grudnia 2009 r.)

III. Uwagi ogólne

1. Konsolidacja regulacji gospodarki danych

Włączenie i skonsolidowanie szeregu regulacji cyfrowych, tj. aktu w sprawie zarządzania danymi, rozporządzenia 2018/1807 w sprawie ram swobodnego przepływu danych nieosobowych, rozporządzenia 2019/1150 w sprawie propagowania sprawiedliwości i przejrzystości dla użytkowników biznesowych korzystających z usług pośrednictwa internetowego oraz dyrektywy 2019/1024 w sprawie otwartych danych i ponownego wykorzystywania informacji sektora publicznego w akcie w sprawie danych stanowi krok spóźniony, lecz kierunkowo właściwy. Obecny stan prawny charakteryzujący się fragmentacją instrumentów regulacyjnych bez wzajemnych odniesień normatywnych jest trudny do uzasadnienia z perspektywy skuteczności i spójności systemowej prawa UE.

2. Teza centralna dotycząca fragmentacji regulacyjnej

AI Act sam w sobie nie stanowi głównej przeszkody dla przedsiębiorstw. Największym problemem pozostaje fragmentacja regulacyjna i nakładanie się instrumentów unijnych.

3. Ogólny cel uproszczenia

Pozytywnie należy ocenić ogólny cel regulacji zarówno Digital Omnibus on AI, jak i Digital Omnibus jakim jest uproszczenie regulacji i zniesienie niektórych barier administracyjnych. Zgodzić się należy z EDPB i EDPS, które we wspólnej opinii 1/2026 wspierają dążenie do ułatwienia wdrożenia AI Act i zmniejszenia obciążeń administracyjnych, jednakże pod warunkiem, że uproszczenia te nie prowadzą do obniżenia poziomu ochrony praw podstawowych osób fizycznych, w szczególności prawa do ochrony danych osobowych. Pozytywnie należy oceniać takie rozwiązania jak wprowadzenie pojedynczego punktu kontaktowego, regulacje umożliwiające opracowanie wspólnego wzoru i wspólnej metodyki przeprowadzania ocen skutków dla ochrony danych, czy zmiany w regulacji dotyczącej cookies.

4. Współpraca AI Office z organami krajowymi

Pozytywnie należy ocenić wymóg aktywnej współpracy między AI Office a krajowymi organami właściwymi przy nadzorowaniu systemów AI opartych na modelach GPAI. Centralizacja nadzoru nad określonymi kategoriami systemów AI może przyczynić się do spójności egzekwowania przepisów AIA na poziomie UE.

5. Krytyka społeczna i instytucjonalna

O ile zasadniczy cel deregulacyjny należy ocenić pozytywnie, to jednak znaczna część szczegółowych proponowanych zmian budzi poważne zastrzeżenia z perspektywy poziomu ochrony praw i wolności. Odnotować należy, że propozycje zawarte w pakietach Digital Omnibus spotkały się z bezprecedensową krytyką. Ponad 120 organizacji społeczeństwa obywatelskiego wystosowało list otwarty, określając propozycje jako "największy rollback praw cyfrowych w historii UE". Krytyka obejmuje również organizacje takie jak European Digital Rights (EDRI), NYOB, a także instytucje akademickie i think-tanki.

IV. Uwagi Szczegółowe

1. Zmiany w RODO (Rozporządzenie 2016/679)

1.1. Uwagi szczegółowe

1.1.1. Ocena ogólnego kierunku zmian

Proponowane zmiany w RODO, mimo deklarowanego celu uproszczenia regulacji i zwiększenia konkurencyjności europejskiej gospodarki cyfrowej, budzą poważne wątpliwości z perspektywy ochrony praw podstawowych i spójności systemu prawnego Unii Europejskiej. Dążenie do ułatwienia wdrożenia AI i zmniejszenia obciążeń administracyjnych zasługuje na wsparcie, jednakże wyłącznie pod warunkiem, że uproszczenia te nie prowadzą do obniżenia poziomu ochrony praw podstawowych osób fizycznych, w szczególności prawa do ochrony danych osobowych zagwarantowanego w art. 8 Karty Praw Podstawowych UE.

1.1.2. Naruszenie zasady neutralności technologicznej

Proponowany art. 88c RODO, który miałby uznawać rozwój systemów AI za prawnie uzasadniony interes w rozumieniu art. 6 ust. 1 lit. f) RODO, budzi fundamentalne zastrzeżenia z perspektywy zasady neutralności technologicznej. Dotychczas art. 6 ust. 1 RODO był technologicznie neutralny – po raz pierwszy konkretna technologia (nie cel przetwarzania) byłaby swoiście legitymizowana w przepisie. Takie podejście może implikować, że przetwarzanie jest legalne już dlatego, że wykorzystywana jest technologia AI – podczas gdy w innych przypadkach to samo przetwarzanie nie mieściłoby się w zakresie art. 6 ust. 1 RODO. Stanowi to naruszenie fundamentalnej zasady, na której opiera się europejskie prawo ochrony danych. Analogicznie zasadę neutralności technologicznej regulacji narusza proponowana zmiana art. 9 ust. 2 w zakresie odnoszącym się do dodania lit. k.

1.1.3. Zbędność interwencji legislacyjnej

Stanowisko dotyczące możliwości oparcia rozwoju systemów AI na prawnie uzasadnionym interesie jest zasadniczo zgodne z tym, co zostało już wyartykułowane w opinii EDPB 28/2024 w sprawie przetwarzania danych osobowych w kontekście modeli

AI. Opinia ta uznaje, że rozwój AI może co do zasady opierać się na prawie uzasadnionym interesie przy zachowaniu ścisłych zabezpieczeń. Skoro zatem istnieje już odpowiednia wykładnia prawa, zasadne jest pytanie o celowość interwencji legislacyjnej. Właściwe stosowanie istniejących mechanizmów przez organy nadzorcze i EDPB wydaje się wystarczające, bez konieczności nowelizacji materialnych przepisów RODO. Interpretacja prawa leży bowiem w gestii organów stosujących prawo, nie zaś legislatora. Analogiczne zastrzeżenia budzi również m.in. proponowana zmiana odnosząca się do pojęcia danych osobowych, która jest naszym zdaniem nie tylko zbędna, ale i istotnie zawęża zakres zastosowania RODO, o czym szerzej w części szczegółowej.

1.1.4. Alternatywne podejście regulacyjne

RODO zapewnia adekwatne i elastyczne mechanizmy regulacyjne, przede wszystkim poprzez szeroką pojemność normatywną zasad przetwarzania danych oraz proceduralne podejście oparte na ryzyku. Większy potencjał dla skutecznej ochrony praw podstawowych przy jednoczesnym wspieraniu innowacji leży nie w zmianach legislacyjnych RODO, lecz w działaniach alternatywnych. W perspektywie krótkoterminowej należałoby rozważyć:

- opracowanie wytycznych EDPB na podstawie art. 70 ust. 1 lit. e) RODO dotyczących kryteriów przejrzystości AI, prawa sprzeciwu wobec przetwarzania danych treningowych, ograniczenia celu w kontekście szkolenia AI oraz odpowiednich zabezpieczeń w rozumieniu art. 10 ust. 5 AIA;
- tworzenie piaskownic regulacyjnych na podstawie art. 57 ust. 1 i 2 AIA;
- wzmocnienie organów nadzorczych poprzez utworzenie wyspecjalizowanych struktur ds. AI; oraz
- skuteczne egzekwowanie rozporządzenia 2025/2518 z dnia 26 listopada 2025 r. w sprawie ustanowienia dodatkowych przepisów proceduralnych dotyczących egzekwowania rozporządzenia (UE) 2016/679.

1.2. Uwagi szczegółowe

1.2.1. Zmiana definicji danych osobowych (art. 4 ust. 1 RODO)

a) Zniekształcenie logiki motywu 26 RODO

Projekt zmienia logikę motywu 26 – przenosząc jego fragmenty do przepisów, ale zniekształcając sens i kontekst. Motyw 26 RODO stanowił wyjątek interpretacyjny, precyzyjnie ograniczony do sytuacji, w których osoba nie jest identyfikowalna w sposób rozsądnie prawdopodobny, przy uwzględnieniu wszystkich podmiotów mogących racjonalnie dysponować środkami identyfikacji. Projekt przenosi do art. 4 ust. 1 jedynie „wycięty fragment” motywu, ale pomija element oceny wszystkich rozsądnie prawdopodobnych środków, pomija wymóg oceny kontekstu technologicznego, ekonomicznego i infrastrukturalnego oraz pomija ryzyko identyfikacji przez inne podmioty w ekosystemie danych. Identyfikacja nie jest bowiem wyłącznie cechą danych, lecz cechą ekosystemu, którego projekt nie bierze pod uwagę – co prowadzi do błędnego częściowego zawężenia definicji danych osobowych.

b) Deformacja sensu motywu 26

Projekt deformuje sens motywu 26 – zastępując analizę ryzyka identyfikacji analizą administracyjną ograniczoną do jednego podmiotu. Motyw 26 wymaga analizy wszystkich rozsądnie dostępnych środków, wszystkich stron mogących wejść w posiadanie danych, kosztów i nakładów, dostępnych technologii oraz postępu algorytmicznego. Projekt Komisji sprowadza tę analizę do pytania: „czy ten konkretny podmiot może zidentyfikować osobę?”. To fundamentalna zmiana prowadząca do skrajnej relatywizacji danych osobowych do zasobu administratora, a nie do rzeczywistego ryzyka ingerencji w prawa i wolności osób fizycznych i kryteriów kontekstowo zobiektywizowanych odnoszących się do ekosystemu. Przesuwa to dalej charakter pojęcia danych osobowych w stronę jego subiektywizacji, co istotnie zawęża jego zakres.

c) Niekontrolowana erozja definicji

Przeniesienie fragmentu motywu 26 do artykułu 4 tworzy niekontrolowaną erozję definicji danych osobowych. Motywy są narzędziem interpretacyjnym, nie normatywnym. Włączenie ich do artykułu prawotwórczego czyni wyjątek regułą, podnosi interpretację do rangi normy oraz zmienia strukturę ochrony danych wbrew konstrukcji RODO. Motywy mają rolę dynamiczną, dostosowującą się do dorobku TSUE. Zastąpienie

ich przepisem blokuje tę elastyczność i prowadzi do przeregulowania w kierunku korzystnym dla administratorów.

d) Ignorowanie dorobku TSUE

Zmiana ignoruje dorobek TSUE, którego motyw 26 jest efektem. TSUE w orzeczeniach *Breyer* (C-582/14), *Nowak* (C-434/16), *Jehovan todistajat* (C-25/17) precyzyjnie zbudował kryterium identyfikowalności relatywnej, uwzględniając elementy kontekstowe, co ostatecznie znalazło odzwierciedlenie w motywie 26. Projekt Komisji ogranicza identyfikowalność do perspektywy jednego podmiotu, zrywa spójność między motywem 26 a art. 4 oraz tworzy nową, formalistyczną definicję, która odchodzi od „materialnego testu ryzyka” stosowanego przez TSUE. To osłabia rolę dorobku sądowego i prowadzi do „zamrożenia” przepisów w momencie, w którym algorytmy identyfikacji dynamicznie się rozwijają.

e) Ograniczenie zakresu stosowania RODO sprzeczne z jego celami

Zmiana ma potencjał ograniczenia zakresu stosowania RODO w sposób sprzeczny z jego celami (art. 1, motyw 4). Definicja danych osobowych jest bramą wejściową dla stosowania całego systemu ochrony danych. Proponowana zmiana istotnie ogranicza pojęcie danych osobowych wbrew celom rozporządzenia i orzecznictwu TSUE, apriorycznie ograniczając podmiotowo relatywność. Jest to sprzeczne z zasadniczą kontekstowością regulacji oraz sprzeczne z podejściem opartym na ryzyku. W konsekwencji zmiana zawęży zakres stosowania RODO, pozwala administratorom twierdzić, że „nie dysponują środkami identyfikacji” oraz ułatwia komercyjne grupowanie, profilowanie i analizę danych rzekomo „nieosobowych”. W środowisku AI i *big data* takie dane mogą być automatycznie odtworzone, zrekonstruowane lub skorelowane z innymi zbiorami, przez co ryzyko identyfikacji pozostaje realne.

f) Faworyzowanie interesów rynkowych kosztem ochrony jednostki

Zmiana sprzyja administratorom i interesom rynkowym kosztem ochrony jednostki. Projekt wprost ułatwia: korzystanie z danych pseudonimizowanych jako rzekomo nieosobowych, przekazywanie danych do podmiotów trzecich bez obowiązków RODO, wyłączenie obowiązków informacyjnych wobec osób, których dane dotyczą oraz wykonywanie analiz i trenowanie modeli AI bez kontroli. Wydaje się to być efektem presji biznesowej i politycznej, a nie potrzeby prawnej. Motyw 26 był instrumentem ochronnym, a nie narzędziem deregulacji.

g) Sprzeczność z zasadą ostrożności wobec ryzyk AI

Zmiana stawia RODO w sprzeczności z zasadą ostrożności wobec ryzyk AI. W kontekście AI definicja danych osobowych powinna być interpretowana szeroko i *pro-persona*, ponieważ: dane pierwotnie „anonimowe” mogą być łatwo reidentyfikowane, dane agregowane mogą zawierać informacje o jednostkach, identyfikacja może być probabilistyczna i algorytmiczna, a ekosystem danych jest globalny, a nie „administrator-silosowy”. Projektowana zmiana ignoruje te procesy i utrwała definicję nieadekwatną do ryzyk AI.

h) Brak oceny skutków dla ochrony praw i wolności człowieka

Zmiana pozbawiona jest dogłębnej oceny skutków dla ochrony praw i wolności człowieka. Komisja nie przedstawiła analizy wpływu na prawa podstawowe (art. 7–8 KPP), oceny ryzyka systemowego identyfikacji w AI ani oceny wpływu na procesy decyzyjne w administracji publicznej, finansach, zdrowiu i edukacji. Brak takich analiz jest szczególnie niepokojący, biorąc pod uwagę, że motyw 26 dotyczył analizy ryzyka, a projekt prawie całkowicie usuwa element ocenny i kontekstowy.

1.2.2. Zmiany zasady ograniczenia celu przetwarzania (art. 5 ust. 1 lit. b RODO)

Choć art. 5 ust. 1 lit. b RODO w obecnym brzmieniu przewiduje domniemanie zgodności dalszego przetwarzania danych do celów naukowych, to ochrona praw osób, której wyrazem jest również motyw 159 RODO, wymaga każdorazowej analizy proporcjonalności, konieczności oraz ryzyka, dokonywanej w ramach art. 89 ust. 1. Proponowana zmiana — poprzez wyłączenie stosowania art. 6 ust. 4 — usuwa istotny element tej analizy, upraszczając proces kosztem poziomu ochrony danych i spójności systemowej. Nie wynika ona z logiki RODO ani z dorobku TSUE, lecz stanowi zabieg deregulacyjny mający charakter polityczno-gospodarczy, co wynika także z motywu 29 projektu.

a) Ryzyko wyłączenia testu zgodności celów

Proponowana zmiana, podobnie jak zmiana poprzez dodanie definicji „*scientific research*”, prowadzić może do ryzyka braku obowiązku realizowania testu, o którym mowa w art. 6 ust. 4 RODO. Tymczasem test ten stanowi istotny element systemu ochrony, zapewniający, że dalsze przetwarzanie danych pozostaje zgodne z pierwotnymi celami ich zebrania lub jest z nimi kompatybilne w świetle uzasadnionych oczekiwań osób, których dane dotyczą.

b) Problem braku definicji legalnej badań naukowych

W RODO brak definicji legalnej pojęcia badań naukowych. Jednocześnie, odwołując się do motywu 159 RODO, należy podkreślić, że pojęcie badań naukowych musi uwzględniać jego powszechne znaczenie (*common meaning and understanding*). W związku z tym przez „badania naukowe” należy rozumieć projekt badawczy ustanowiony zgodnie z odpowiednimi sektorowymi standardami metodologicznymi i etycznymi, w zgodności z dobrą praktyką. Pojęcie to obejmuje systematyczną aktywność, w tym gromadzenie i analizę danych, która zwiększa zasób zrozumienia i wiedzy oraz ich zastosowanie. Co istotne, nie wyklucza się z pojęcia badań naukowych badań komercyjnych, jednak rola badań jest rozumiana jako dostarczanie wiedzy, która z kolei może poprawić jakość życia wielu osób i poprawić wydajność usług socjalnych.

c) Ograniczenie dalszego wykorzystania wyników badań

Dalsze wykorzystanie danych powinno być możliwe, jeśli dotyczy wyników badań, a nie danych, na podstawie których badania były realizowane. Proponowana zmiana nie uwzględnia tego fundamentalnego rozróżnienia, co może prowadzić do nieograniczonego wtórnego wykorzystania surowych danych osobowych pod pretekstem „badań naukowych”.

d) Unormatywnienie motywu 159 RODO

Wprowadzenie definicji pojęcia badań naukowych, podobnie jak przeniesienie częściowe treści motywu 26 do art. 4 pkt 1, stanowi unormatywnienie motywu 159 RODO. Ponownie podkreślić należy, że motywy są narzędziem interpretacyjnym, nie normatywnym. Włączenie ich do artykułu prawotwórczego powoduje, że podnosi interpretację do rangi normy, ogranicza się elastyczność kontekstową i dynamiczność znaczeniową, co w kontekście neutralności technologicznej regulacji RODO ma bardzo istotne znaczenie, zwłaszcza uwzględniając orzecznictwo TSUE. Zastąpienie motywów przepisem blokuje tę elastyczność i prowadzi do przeregulowania w kierunku korzystnym dla administratorów. Znaczenie motywów ma silne powiązanie kontekstowe, a użyte pojęcia poddają się wykładni przez pryzmat zasad i dyrektyw regulacyjnych UE. W kontekście UE są to zwłaszcza zasady godnej zaufania sztucznej inteligencji, które stanowią punkt odniesienia i kontekstowo warunkują wykładnię, również przepisów RODO, w tym dotyczących pojęcia badań naukowych, przy uwzględnieniu celów ochronnych RODO.

e) Rozszerzenie pojęcia badań naukowych o cele komercyjne

Pojęcie badań naukowych zostaje wbrew dotychczasowym intencjom prawodawcy wyrażonym w motywie 159 rozszerzone o działalność prowadzoną w istocie w celu komercyjnym, co z uwagi na specyfikę cyklu życia AI, w istocie prowadzi do tego, że pojęciem badań naukowych można objąć fazę trenowania każdego systemu AI. To fundamentalne wypaczenie sensu wyjątku badawczego, który miał służyć rozwojowi wiedzy w interesie publicznym, a nie komercyjnemu wykorzystaniu danych osobowych.

1.2.3. Przetwarzanie szczególnych kategorii danych w kontekście AI (art. 9 ust. 2 lit. k) i ust. 5 RODO)

Konstrukcja proponowanego art. 9 ust. 2 lit. k) oraz art. 9 ust. 5 RODO dotycząca przetwarzania danych wrażliwych w kontekście rozwoju i operacji systemów AI jest wadliwa z wielu powodów systemowych i praktycznych.

a) Wywrócenie systemu art. 9 RODO

Wprowadzenie nowej podstawy prawnej przetwarzania danych wrażliwych wywraca system art. 9 RODO. Art. 9 RODO ma konstrukcję wybitnie wyjątkową: zasadą jest całkowity zakaz przetwarzania danych szczególnych kategorii, wyjątkiem są taksatywnie wyliczone, ściśle określone przesłanki, a każdy wyjątek musi być interpretowany ściśle i zawężająco. Dodanie lit. k („AI system or AI model”) rozszerza katalog wyjątków w sposób niespójny systemowo, oderwany od logiki ochrony danych osobowych szczególnych kategorii oraz zorientowany na interes technologiczny, a nie ochronny. Jak wskazuje TSUE, każda ingerencja musi odpowiadać konieczności, proporcjonalności, by nie ingerować w prawa podstawowe w sposób ogólny i nieodróżnicowany – czego projekt nie zapewnia (m.in. C-293/12 i C-594/12, *Digital Rights Ireland*, pkt 52–54)

b) Ekstremalnie szeroki zakres wyjątku

Zakres wyjątku jest ekstremalnie szeroki i obejmuje praktycznie wszystkie etapy życia systemów AI. Przepis obejmuje trenowanie, testowanie, walidację oraz działanie (*operation*) modeli AI. Oznacza to, że wyjątek nie dotyczy „konkretnej czynności”, lecz całego cyklu życia modelu AI, a każde przetwarzanie danych wrażliwych w ramach AI – również komercyjnego, masowego, nieweryfikowalnego – może być uznane za zgodne z prawem. To skrajnie niebezpieczne, ponieważ żaden inny wyjątek z art. 9 ust. 2 nie ma takiej ogólnej struktury.

c) Brak wymogu niezbędności i proporcjonalności

Przepis nie wymaga wykazania niezbędności ani proporcjonalności. W lit. k brakuje dwóch zasadniczych przesłanek ochronnych, obecnych w większości innych wyjątków: „niezbędności” przetwarzania oraz „proporcjonalności” w świetle celu. Większość wyjątków w art. 9 ust. 2 dotyczy zdrowia publicznego, ratowania życia, zatrudnienia, ochrony ważnego interesu publicznego lub celów archiwalnych w interesie publicznym, do celów badań naukowych lub historycznych lub do celów statystycznych. Tymczasem lit. k dopuszcza przetwarzanie danych wrażliwych z uwagi na interes techniczno-rozwojowy prywatnego administratora. To fundamentalnie zaburza aksjologię RODO.

d) Brak określenia celu uzasadniającego przetwarzanie

Art. 9 ust. 2 lit. k nie określa, jaki cel uzasadnia przetwarzanie danych wrażliwych. Przesłanki art. 9 ust. 2 są ściśle powiązane z ochroną zdrowia, prawem pracy, ważnym interesem publicznym itp. Lit. k nie wskazuje żadnego konkretnego, materialnego celu: „*processing in the context of the development and operation of an AI system*”. Tak szerokie sformułowanie oznacza legalizację prawie każdego modelu AI, brak konieczności wykazania celu publicznego oraz brak rozróżnienia pomiędzy interesem prywatnym a interesem publicznym. Oznacza to, że podmioty rozwijające technologie komercyjne będą mogły przetwarzać dane wrażliwe dla swoich modeli tak łatwo, jak jednostka badawcza w interesie publicznym — co jest sprzeczne z aksjologią RODO.

e) Techniczna nierealistyczność art. 9 ust. 5

Przepis zakłada usuwanie danych wrażliwych „jeśli zostaną wykryte” — co jest technicznie nierealne. Art. 9 ust. 5 wymaga identyfikacji danych wrażliwych w datasetach, a następnie ich usunięcia albo ochrony. Problemy są następujące: AI nie jest w stanie wykryć wszystkich danych wrażliwych, zwłaszcza inferowanych; w modelach ML dane wrażliwe mogą być zaszyte w parametrach i być nieusuwalne; usunięcie danych nie usuwa informacji zawartych w wagach; administrator nie ma narzędzi, by ocenić, czy model „nadal pamięta” dane wrażliwe; mechanizmy „*machine unlearning*” są w fazie eksperymentalnej. W konsekwencji przepis opiera się na fikcji technicznej.

f) Furtka do omijania zgody

Wyjątek lit. k będzie nadużywany jako furtka do omijania zgody z art. 9 ust. 2 lit. a. W kontekście AI większość przetwarzania danych wrażliwych powinna odbywać się na

podstawie szczególnej zgody (art. 9 ust. 2 lit. a) lub w ramach przetwarzania publicznego (lit. g–j). Lit. k umożliwi administratorom rezygnację z uzyskiwania zgody lub spełniania innych podstaw przetwarzania danych szczególnych kategorii.

g) Brak rozróżnienia faz *development* i *operation*

Przepis nie wprowadza szczegółowego rozróżnia przetwarzania eksploracyjnego (*development*) i operacyjnego (*operation*) obejmując tworząc podstawy przetwarzania jednakowo dla obu faz, pomijając różnice w potencjalnym poziomie oddziaływania przetwarzania w poszczególnych na prawa i wolności podmiotów danych.

h) Właściwe miejsce regulacji

Właściwym miejscem regulacji przetwarzania szczególnych kategorii danych dla celów wykrywania i korygowania uprzedzeń (*bias detection*) jest art. 10 ust. 5 AIA — i tam należy ją pozostawić, przy ewentualnym rozszerzeniu zakresu podmiotowego i przedmiotowego w sposób *ściśły i ograniczony*. Zasadne jest utrzymanie standardu „ściśłej konieczności” (*strict necessity*), obecnie obowiązującego dla systemów wysokiego ryzyka, dla wszystkich dostawców i podmiotów wdrażających systemy AI w zakresie przetwarzania szczególnych kategorii danych osobowych, pomimo różnicy zakresów obowiązywania obu tych aktów.

1.2.4. Obowiązki informacyjne (art. 13 RODO)

Propozycja zwolnienia administratorów z pełnych obowiązków informacyjnych, gdy uznają oni, że relacja z użytkownikiem jest „jasna i ograniczona” lub gdy sądzą, że osoba „już posiada” informacje, budzi poważne zastrzeżenia.

Należy podkreślić, że w obowiązującym brzmieniu przepisów **art. 13 ust. 4 RODO przewiduje już wyłączenie**, zgodnie z którym ust. 1, 2 i 3 nie mają zastosowania, gdy — i w zakresie, w jakim — osoba, której dane dotyczą, dysponuje już tymi informacjami. Proponowana zmiana idzie jednak znacznie dalej, wprowadzając dodatkowe, uznaniowe przesłanki wyłączenia obowiązku informacyjnego.

W praktyce przejrzystość stałaby się opcjonalna — osoby nie byłyby już jasno informowane o tym, jakie dane są zbierane, w jakim celu i jak długo są przechowywane. Bez tej wiedzy prawa takie jak dostęp, sprzeciw czy usunięcie tracą wszelkie praktyczne znaczenie, ponieważ osoba nie wie, że może i powinna z nich skorzystać. Proponowane rozszerzenie wyłączeń stanowi zatem istotne osłabienie fundamentalnej zasady

przejrzystości, która jest warunkiem *sine qua non* skutecznej ochrony danych osobowych.

1.2.5. Zautomatyzowane podejmowanie decyzji (art. 22 RODO)

Proponowana zmiana art. 22 RODO, polegająca zamiast zakazu warunkowe dopuszczenie oraz usunięciu wymogu „niezbędności” jako warunku dopuszczalności zautomatyzowanego podejmowania decyzji w kontekście umownym, stanowi fundamentalną zmianę filozofii regulacji, która miała być najsilniejszym mechanizmem ochronnym gwarantującym autonomię i kontrolę podmiotom danych.

a) Zmiana filozofii ochrony danych — od zakazu do dopuszczenia

Art. 22 w obowiązującym brzmieniu zbudowany na idei ochrony przed automatyzacją, ochrony autonomii jednostki, gwarancji minimum ludzkiej kontroli oraz zasady ostrożności wobec zautomatyzowanego podejmowania decyzji. Digital Omnibus zmienia to w kierunku normalizacji zautomatyzowanego podejmowania decyzji, umożliwienia masowego stosowania automatyzacji oraz faktycznej rezygnacji z decyzji podejmowanej przez człowieka, na etapie przed podjęciem decyzji.

b) Praktyczne znaczenie usunięcia wymogu niezbędności

Dotychczas zautomatyzowane podejmowanie decyzji można było w relacjach kontraktowych stosować tylko gdy bez niego realizacja umowy była niemożliwa. Przykładowo: *scoring* kredytowy oparty wyłącznie na zautomatyzowanym podejmowaniu decyzji mógł być stosowany tylko, jeśli bank nie mógł zapewnić alternatywy. Po zmianie administrator może samodzielnie zdecydować, że stosuje zautomatyzowane podejmowanie decyzji, bez konieczności wykazania, że jest to niezbędne do zawarcia lub wykonania umowy. Konsekwencją jest to, że zautomatyzowane podejmowanie decyzji staje się opcją biznesową administratora, a nie wyjątkiem chroniącym osobę fizyczną. To jest całkowite odwrócenie dotychczasowej logiki.

c) Zautomatyzowane podejmowanie decyzji jako standard biznesowy

W wielu sektorach (bankowość, ubezpieczenia, telekomunikacja, praca, zdrowie, edukacja, media) administratorzy będą mogli stosować zautomatyzowane podejmowanie decyzji jako *default*, argumentować, że jest to „efektywne”, „tańsze”, „szybsze”, eliminować procedury manualne oraz przerzucać decyzje o dużej wadze na algorytmy. Zmiana regulacyjna zatem jest zmianą paradygmatu.

d) Utrata możliwości realnego wyboru przez osobę fizyczną

Bez wymogu niezbędności osoba nie może żądać alternatywy, administrator może odmówić świadczenia usługi, jeśli odmówi się zautomatyzowanego podejmowania decyzji, a praktyka rynkowa wymusi zgodę na zautomatyzowane podejmowanie decyzji jako warunek korzystania z większości usług. To prowadzi do algorytmicznego przymusu ekonomicznego.

e) Iluzoryczność prawa do interwencji ludzkiej

Prawo do interwencji ludzkiej (art. 22 ust. 3) staje się iluzoryczne. Jeżeli zautomatyzowane podejmowanie decyzji jest dopuszczalne bez ograniczeń, to interwencja ludzka staje się *post factum*, proces decyzyjny jest w pełni automatyczny, a rola człowieka ograniczy się do „wyjaśnienia decyzji”, a nie jej podjęcia, co jest podejściem zasadniczym zgodnie z obecnie obowiązującymi przepisami.

f) Kolizja z orzecnictwem TSUE

TSUE w sprawach C-634/21 *SCHUFA*, C-184/20 *Vyriausioji tarnybinės etikos komisija* podkreślał, że odstępstwa i ograniczenia zasady ochrony takich danych powinny być stosowane w granicach tego, co jest ściśle konieczne, zautomatyzowane podejmowanie decyzji musi być ściśle ograniczone, test niezbędności jest decydujący oraz wyrównanie asymetrii między administratorem a osobą fizyczną ma charakter konstytutywny. Digital Omnibus ignoruje tę linię orzecniczą.

1.2.6. Prawnie uzasadniony interes dla systemów AI (art. 88c RODO)

Proponowany art. 88c RODO rodzi szereg problemów praktycznych i systemowych, które wykraczają poza kwestię neutralności technologicznej.

a) Systemowe osłabienie ochrony danych osobowych

Art. 88c tworzy sektorowy wyjątek dla technologii AI, który rozszerza zakres dopuszczalnego przetwarzania danych (w tym wtórnego), uprzywilejowuje interes administratora nad prawami jednostki, odrywa interes administratora od wymaganego testu niezbędności i prawnego uzasadnienia, sprowadzając do deklaracji istnienia oraz czyni z podstawy prawnie uzasadnionego interesu mechanizm pierwszego wyboru,

czyniąc z praw i wolności jednostki wyjątek a nie cel ochrony. To odwraca aksjologię RODO, którego celem nadrzędnym jest ochrona praw podstawowych, nie wspieranie rozwoju technologii, przy braku hierarchii podstaw prawnych przetwarzania. Wprowadza w neutralnej technologicznie regulacji wyłom o charakterze ściśle technologicznym.

b) Sprzeczność z opinią EROD 28/2024

Opinia EROD 28/2024 nie daje podstaw do uprzywilejowania LIA dla AI. Przeciwnie, podkreśla: brak hierarchii podstaw prawnych, konieczność pełnego testu LIA (interes – konieczność – równowaga), szczególne ryzyka związane z działaniem modeli AI oraz brak rozsądnych oczekiwań osób fizycznych co do wykorzystywania ich danych do trenowania modeli AI. Art. 88c ignoruje te wymogi, tworząc „ustawowe uproszczenie” sprzeczne z wytycznymi EROD.

c) Kolizja z orzecznictwem TSUE

TSUE wymaga ścisłej wykładni wyjątków, oceny kontekstu i asymetrii informacji oraz szczególnej ochrony wobec procesów profilowania. Art. 88c dopuszcza przetwarzanie danych w fazie rozwoju i działania modeli AI bez testu niezbędności, bez analizy zgodności celu oraz w warunkach skrajnej asymetrii wiedzy między administratorem a osobą fizyczną. Jest to sprzeczne z linią orzeczniczą (m.in. *Meta*, *SCHUFA*, *Nowak*, *Jehovan todistajat*). Zgodnie z orzecznictwem TSUE (C-252/21 *Bundeskartellamt*, pkt 112; C-621/22 *Tennisbond*, pkt 55), rozsądne oczekiwania osoby, której dane dotyczą, odgrywają kluczową rolę w teście równowagi interesów. W przypadku szkolenia i operacji AI złożoność, wielość i ciągła ewolucja systemów oznaczają, że osoby, których dane dotyczą, nie mogą rozsądnie oczekiwać takiego przetwarzania ani jego zakresu.

d) Naruszenie kluczowych zasad RODO

W zakresie *ograniczenia celu* (art. 5 ust. 1 lit. b) – art. 88c umożliwia wtórne przetwarzanie danych osobowych do celów AI bez testu zgodności celów, co związane jest z wyjęciem go poza nawias regulacji podstaw prawnych jedynie z odwołaniem do art. 6 ust. 1 lit. f. W zakresie *minimalizacji danych* (art. 5 ust. 1 lit. c) – modele AI wymagają danych masowych i zróżnicowanych; art. 88c nie zapewnia realnych mechanizmów minimalizacyjnych. W zakresie *rozliczalności* (art. 5 ust. 2) – przeniesienie ciężaru oceny legalności na samą deklarację istnienia „interesu” administratora – bez wymogu wykazania niezbędności czy proporcjonalności – osłabia rozliczalność. W zakresie *prawa do sprzeciwu* (art. 21) – w praktyce jest nieskuteczne: dane wykorzystane do trenowania modeli nie podlegają odwracalnemu usunięciu.

2. Zmiany w AIA (Rozporządzenie 2024/1689)

2.1. Uwagi ogólne

2.1.1. Ocena ogólnego kierunku zmian

Deklarowanymi celami propozycji Digital Omnibus on AI są m.in. uproszczenie procesu stosowania AIA czy zmniejszenie obciążeń administracyjnych dla przedsiębiorstw. Cele te jako takie zasługują na aprobatę. Warto jednak mieć na względzie, że w praktyce wiele z proponowanych zmian – wydaje się nie mieć charakteru wyłącznie technicznego, porządkującego czy „doprecyzowującego” istniejące przepisy, a idąc dalej, niektóre z nich mogą stwarzać poważne ryzyko, że będą odbywać się kosztem osłabienia mechanizmów ochrony praw podstawowych, rozliczalności oraz egzekwowalności obowiązków dostawców i podmiotów stosujących czy nawet prowadzić do deregulacji systemowej.

2.1.2. Tryb procedowania projektu

Propozycja Digital Omnibus on AI jest procedowana w wyjątkowo szybkim tempie, co rodzi poważne ryzyko zarówno niepełnej oceny Komisji Europejskiej co do skutków tej regulacji, jak i pominięcia w tym procesie odpowiednio szerokich konsultacji publicznych. Na powyższe wskazują także informacje wynikające z dokumentów Komisji, upublicznionych przez Corporate Europe Observatory (która uzyskała je w trybie dostępu do informacji publicznej). Z materiałów tych wynika, że w konsultacjach poprzedzających propozycję (tzw. „Reality Checks”) udział brali przede wszystkim przedstawiciele „biznesu”, w tym warto zauważyć, że jeśli chodzi o „Reality Check on AI” – uczestniczyło w nim dziesięć przedsiębiorstw i stowarzyszeń branżowych (w tym podmiotów postulujących za „osłabieniem” AIA) i zaledwie dwie organizacje społeczeństwa obywatelskiego.

Biorąc pod uwagę skalę wpływu sztucznej inteligencji na rzeczywistość gospodarczą, prawną i społeczną – to tak istotne zmiany w regulacji tego obszaru – powinny być poprzedzone odpowiednio szerokimi konsultacjami zapewniającymi przede wszystkim zrównoważoną reprezentację wszystkich zainteresowanych stron, w tym administracji publicznej, nauki, organizacji konsumenckich, środowisk praw człowieka, organizacji społeczeństwa obywatelskiego etc.

2.1.3. Wpływ na prawa podstawowe

Proponowane zmiany w AIA (rozpatrywane w ujęciu systemowym) wydają się prowadzić do osłabienia mechanizmów ochrony praw podstawowych przewidzianych w pierwotnym tekście rozporządzenia, a co najmniej stwarzają realne ryzyko takiego osłabienia.

Choć poszczególne proponowane zmiany w AIA mogą w założeniu nieść ze sobą uzasadnione uproszczenia regulacyjne (i zapewne w praktyce tak może być), to brak pełnej oceny skutków, przyspieszony tryb procedowania oraz niedostateczna reprezentacja interesariuszy spoza sektora przemysłu nie pozwalają na rzetelną ocenę rzeczywistego wpływu propozycji na prawa podstawowe i zagrożeń jakie mogą te zmiany spowodować. Analiza poszczególnych zmian w ich wzajemnym powiązaniu wskazuje na ogólne ryzyko systemowego osłabienia mechanizmów ochrony praw podstawowych.

2.2. Uwagi szczegółowe

2.2.1. Zmiana harmonogramu stosowania (art. 111 i 113 AIA)

AIA przewiduje stosowanie przepisów o systemach wysokiego ryzyka od 2 sierpnia 2026 r. (załącznik III – m.in. biometria, wymiar sprawiedliwości, edukacja, zatrudnienie, migracja) i od 2 sierpnia 2027 r. (załącznik I – m.in. wyroby medyczne, lotnictwo, zabawki).

Ponadto art. 111 ust. 2 AIA przewiduje tzw. klauzulę „grandfatheringu”. Zgodnie z tym przepisem do systemów AI wysokiego ryzyka, które zostały wprowadzone do obrotu lub oddane do użytku przed dniem 2 sierpnia 2026 r., co do zasady nie stosuje się obowiązków przewidzianych w rozporządzeniu (co ma zapobiec jego retroaktywnemu stosowaniu). Obowiązki te znajdą zastosowanie dopiero w przypadku, gdy po tej dacie system zostanie poddany istotnym zmianom projektowym (significant changes in design). Jednocześnie przepis ustanawia odrębny termin dostosowawczy dla systemów AI wysokiego ryzyka przeznaczonych do użycia przez organy publiczne – dostawcy i podmioty stosujące powinni podjąć działania zmierzające do zapewnienia zgodności z wymaganiami rozporządzenia najpóźniej do 2 sierpnia 2030 r.

Propozycja Digital Omnibus on AI przewiduje zmianę harmonogramu stosowania regulacji poprzez wprowadzenie tzw. mechanizmu “stop-the-clock” tj. warunkowego odroczenia

stosowania uzależnionego od decyzji Komisji potwierdzającej dostępność narzędzi wspierających zgodność, z datą graniczną nie później niż 2 grudnia 2027 r. (dla systemów z załącznika III) i 2 sierpnia 2028 r. (dla systemów z załącznika I).

Proponowany projekt odraczając stosowanie przepisów rozdziału III w odniesieniu do systemów wysokiego ryzyka z art. 6 ust. 2 i załącznika III, w praktyce wydłuża również okres przejściowy wynikający z klauzuli „grandfatheringu” (art. 111 ust. 2) – data graniczna przesuwana się z 2 sierpnia 2026 r. do maksymalnie 2 grudnia 2027 r. W konsekwencji większa liczba systemów AI wysokiego ryzyka z załącznika III będzie mogła funkcjonować na rynku przez dłuższy czas bez zastosowania pełnego reżimu obowiązków rozdziału III.

a) Wpływ na wymiar sprawiedliwości i administrację publiczną

Odroczenie stosowania przepisów rozdziału III (w zakresie systemów wysokiego ryzyka z załącznika III) oznaczałoby, że nawet do 2 grudnia 2027 r. systemy AI wykorzystywane m.in. w obszarze wymiaru sprawiedliwości, ścigania oraz administracji publicznej (o potencjalnie istotnym wpływie na prawa jednostek) – mogłyby funkcjonować *de facto* bez zastosowania kluczowych wymogów wskazanego rozdziału, w tym: pełnego systemu zarządzania ryzykiem (art. 9), wymogów dotyczących danych i zarządzania danymi (art. 10), obowiązku zapewnienia nadzoru ludzkiego (art. 14), wymogów przejrzystości i przekazywania informacji podmiotom wdrażającym (art. 13) oraz obowiązków podmiotów stosujących, w tym przeprowadzenia oceny wpływu na prawa podstawowe (art. 26–27).

b) Zagrożenie dla ochrony praw podstawowych

Wskazane opóźnienie stosowania kluczowych przepisów rozdziału III oznacza, że potencjalnie przez dodatkowych 16 miesięcy – systemy AI mogące wywierać istotny wpływ na prawa człowieka będą mogły funkcjonować na rynku bez pełnego reżimu ochronnego (tj. bez zastosowania całego pakietu obowiązków i zabezpieczeń jakie przewiduje rozdział III AIA). Dotyczy to w szczególności systemów stosowanych w obszarach biometrii, wymiaru sprawiedliwości, organów ścigania, migracji, edukacji i zatrudnienia, a zatem dziedzin, w których ryzyko naruszenia praw podstawowych jest szczególnie wysokie.

Odroczenie nie oznacza oczywiście w tym okresie całkowitego braku ram prawnych ochrony praw jednostek, ponieważ nadal zastosowanie znajdują m.in. RODO czy standardy ochrony praw podstawowych wynikające z prawa UE, w tym z Karty Praw

Podstawowych. Jednak wydaje się, że regulacje te co do zasady nie adresują w pełni ryzyk specyficznych dla technologii AI, czego potwierdzeniem jest sam fakt uchwalenia AIA jako odrębnego instrumentu prawnego, wprowadzającego dedykowane mechanizmy ochronne wykraczające poza dotychczasowe ramy (m.in. wymogi dotyczące jakości danych treningowych, obowiązek zapewnienia nadzoru ludzkiego czy ocenę wpływu na prawa podstawowe specyficzną dla systemów AI).

c) Kontekst braku standardów technicznych i “stop-the-clock”

Projekt uzasadnia wprowadzenie mechanizmu “stop-the-clock”/“ruchomy termin” – w zakresie stosowania wymogów dla systemów wysokiego ryzyka m.in. tym, że opóźnienia w przygotowaniu instrumentów wspierających zgodność (w szczególności norm zharmonizowanych, wspólnych specyfikacji i wytycznych) mogłyby utrudnić praktyczne wdrożenie obowiązków oraz zwiększyć koszty zgodności. W ocenie Komisji ma to uzasadniać odejście od pierwotnej daty 2 sierpnia 2026 r. Jest to stanowisko uzasadnione, ale jedynie w części.

Z perspektywy ochrony praw podstawowych i pewności prawa należy bowiem podkreślić, że część obowiązków z rozdziału III ma charakter organizacyjny i procesowy (np. system zarządzania ryzykiem, nadzór człowieka, prowadzenie dokumentacji) i może być wdrażana niezależnie od pełnego zestawu norm zharmonizowanych. W konsekwencji samo opóźnienie prac standaryzacyjnych nie powinno automatycznie uzasadniać odroczenia stosowania całego rozdziału III (sekcje 1–3).

Warunkowy i trudny do przewidzenia moment wydania decyzji Komisji utrudnia zapewnienie zgodności po stronie dostawców i podmiotów stosujących systemy AI, a także przygotowanie organów właściwych do nadzoru i egzekwowania przepisów.

d) Postulat różnicowania obowiązków

Zasadne wydaje się rozważenie zróżnicowania odroczenia według rodzaju obowiązków. W szczególności obowiązki, które nie są wprost zależne od standardów technicznych (np. obowiązki transparentności wobec podmiotów stosujących, podstawowe elementy systemu zarządzania ryzykiem, nadzór człowieka, wymogi organizacyjne po stronie podmiotów stosujących w tym przygotowanie do FRIA), mogłyby być stosowane zgodnie z pierwotnym harmonogramem lub z krótszym odroczeniem, przy pozostawieniu ewentualnie dłuższego okresu jedynie dla wymogów stricte technicznych.

e) Graniczne daty (backstop dates)

Pozytywnie należy ocenić wprowadzenie tzw. dat granicznych – 2 grudnia 2027 r. (załącznik III) i 2 sierpnia 2028 r. (załącznik I) – po których przepisy zaczną być stosowane niezależnie od decyzji Komisji. Zapewnia to minimalną przewidywalność co do najpóźniejszego momentu rozpoczęcia stosowania przepisów. Nie eliminuje to jednak niepewności prawnej w okresie poprzedzającym te daty – podmioty zobowiązane nadal nie wiedzą, kiedy dokładnie przepisy zaczną obowiązywać.

2.2.2. Kompetencje w zakresie AI (art. 4 AIA)

Propozycja Komisji zakłada istotną zmianę art. 4 AIA. Obecnie regulacja przewiduje obowiązek wobec dostawców i podmiotów stosujących systemy AI zapewnienia personelowi – w możliwie największym stopniu – odpowiedniej wiedzy, zaś projekt przekształca wskazany obowiązek w zobowiązanie państw członkowskich oraz Komisji do zachęcania tych podmiotów do podejmowania działań na rzecz budowania kompetencji swojego personelu w zakresie AI.

Zniesienie obowiązku i zastąpienie go jedynie “zachęceniem” powoduje przede wszystkim brak realnej możliwości rozliczalności i egzekwowalności standardów AI literacy po stronie dostawców i podmiotów stosujących systemy AI. Nie należy oczekiwać, by standardy AI literacy były w sposób automatyczny spełniane przez adresatów innych obowiązków wynikających z AIA tylko dlatego, że spełnienie tych obowiązków jest funkcjonalnie zależne od posiadania odpowiedniego poziomu kompetencji w zakresie AI. Utrata statusu samodzielnego wymogu w zakresie zgodności z AIA może powodować uznawanie szkoleń w zakresie AI literacy za koszt fakultatywny, w konsekwencji pomijany lub co najmniej nie traktowany priorytetowo.

a) Osłabienie celu regulacji

Zmiana art. 4 AIA znacząco osłabia cel, jakim jest bezpieczne i odpowiedzialne wdrażanie technologii. Brak obowiązku posiadania przez personel “w możliwie największym stopniu” odpowiednich kompetencji – na każdym etapie cyklu życia AI może bezpośrednio przełożyć się na niższy poziom ochrony praw podstawowych, w tym prawa do ochrony danych. Takie stanowisko wyraża również EDPB i EDPS we wspólnej opinii (1/2026).

b) Problem niejasności i proporcjonalności

Zarówno obecne, jak i proponowane brzmienie art. 4 AIA budzi wątpliwości interpretacyjne z uwagi na posługiwanie się pojęciami niedookreślonymi (w szczególności „odpowiedni poziom”), które w praktyce mogą powodować niepewność wdrożeniową co do zakresu wymaganych środków (poziomu merytorycznego i częstotliwości szkoleń czy weryfikacji nabytych kompetencji). Odwołanie się do kryteriów takich jak „wiedza techniczna, doświadczenie, wykształcenie i wyszkolenie” oraz „kontekst wykorzystania” ma charakter ocenny i również nie przesądza o zdefiniowaniu minimalnych standardów. Niepewność pojawia się także co do zakresu podmiotowego tej regulacji tj. ustalenia, którzy pracownicy/współpracownicy faktycznie wymagają przeszkolenia.

c) Utrzymanie szkoleń dla systemów wysokiego ryzyka

Należy zaznaczyć, że zmiana ta nie znosi obowiązków dla podmiotów stosujących systemy AI wysokiego ryzyka. Wymogi w tym zakresie pozostają w mocy, co wynika z brzmienia art. 26 ust. 2, który nakłada na podmioty stosujące te systemy – obowiązek powierzenia nadzoru człowiekowi – osobom fizycznym, które mają niezbędne kompetencje, szkolenie i upoważnienie (oraz zapewnienie im wsparcia). Ten stan normatywny może jednak stanowić podstawę do wzmocnienia przekonania u dostawców i podmiotów stosujących systemy, że standardy AI literacy nie odgrywają istotnej roli w odniesieniu do systemów AI innych niż te wysokiego ryzyka.

d) Ogólna rekomendacja EDPB i EDPS

We wspólnej opinii EDPB i EDPS (1/2026) stanowczo zalecają utrzymanie obowiązkowego charakteru przepisów dotyczących kompetencji w zakresie AI. Organy postulują także zobowiązanie Komisji Europejskiej (i organów regulacyjnych) do wydania dla dostawców i podmiotów stosujących AI, praktycznych wytycznych w zakresie wdrażania AI literacy, zamiast znoszenia obecnego obowiązku wynikającego z art. 4 AIA.

e) Wątpliwości związane z uzasadnieniem zmiany

Uzasadnienie zmiany art. 4, zawartej w motywie 5, budzi wątpliwości, ponieważ opiera się na założeniu o „jednolitym” charakterze pierwotnego obowiązku. Obowiązek zapewnienia AI literacy nie implikował bowiem jednolitego, sztywnego standardu dla wszystkich

podmiotów, lecz – zgodnie z jego funkcją i zasadą proporcjonalności – dopuszczał i zakładał zróżnicowanie zakresu oraz intensywności działań szkoleniowych w zależności od skali działalności, rodzaju systemów i kontekstu ich użycia. Co więcej, nowelizacja nie usuwa wskazywanego problemu, ponieważ także model oparty na „zachęcaniu” pozostaje normatywnie niedookreślony i nie wprowadza żadnych nowych mechanizmów realnego różnicowania adresatów. W konsekwencji argument o „one-size-fits-all” pełni raczej funkcję retorycznego uzasadnienia osłabienia normy niż rzeczywistej diagnozy jej wady konstrukcyjnej.

2.2.3. Zniesienie obowiązku rejestracji systemów AI (zmiana art. 6 ust. 4 oraz uchylenie art. 49 ust. 2 i sekcji B w załączniku VIII AIA)

AIA pozwala dostawcom systemów wymienionych w załączniku III na samodzielną ocenę, w przedmiocie tego czy ich system nie stanowi systemu wysokiego ryzyka, o ile spełnia on przynajmniej jeden z warunków określonych w art. 6 ust. 3 (np. służy wąsko określonej funkcji proceduralnej czy poprawia wynik zakończonych uprzednio czynności wykonywanych przez człowieka) oraz nie stwarza znaczącego ryzyka szkody dla zdrowia, bezpieczeństwa lub praw podstawowych osób fizycznych. Jako swego rodzaju mechanizm równoważący taką swobodę samooceny, AIA nakłada jednocześnie obowiązek rejestracji tych systemów w unijnej bazie danych (art. 49 ust. 2, sekcja B załącznika VIII). Digital Omnibus on AI proponuje zniesienie tego obowiązku, pozostawiając jedynie wymóg udokumentowania oceny przed wprowadzeniem systemu do obrotu lub oddaniem do użytku i udostępnienia jej organom krajowym na ich żądanie.

Zniesienie obowiązku rejestracji może przede wszystkim ograniczyć transparentność samooceny dostawców na etapie *ex ante*. W konsekwencji rzetelne *due diligence* po stronie podmiotów wdrażających oraz możliwość wczesnej reakcji organów właściwych mogą w większym stopniu zależeć od działań podejmowanych dopiero *ex post*, na podstawie dokumentacji udostępnianej na żądanie.

a) Likwidacja mechanizmu kontroli publicznej

Należy podkreślić, że obowiązek rejestracji stanowi jedyną formę uprzedniego (tj. przed wprowadzeniem na rynek) nadzoru nad decyzją dostawcy o uznaniu systemu z załącznika III za niebędący systemem wysokiego ryzyka. Publiczna rejestracja umożliwia podmiotom stosującym AI przeprowadzenie odpowiedniej weryfikacji i oceny ryzyk przed wdrożeniem systemu, a organom krajowym i organom ochrony praw podstawowych (FRABs) – podjęcie działań nadzorczych zanim system trafi na rynek.

EDPB i EDPS w opinii 1/2026 wprost rekomendują utrzymanie tego obowiązku, wskazując że jego zniesienie znacząco osłabi rozliczalność dostawców i stworzy „zachętę” do nadużywania wyłączenia. Zaznaczając także, że oszczędności po stronie dostawców są przy tym znikome w porównaniu z potencjalnymi ryzykami dla praw podstawowych.

b) Przejrzystość i rozliczalność

Ograniczenie obowiązku rejestracji redukuje obciążenia administracyjne dla dostawców systemów, ale istotnie osłabia publiczną widoczność decyzji klasyfikacyjnej, przenosząc kluczowe uzasadnienie poza rejestr – do dokumentacji wewnętrznej dostawcy. W efekcie rozliczalność przyjmuje w większym stopniu charakter warunkowy i reaktywny, zależny od inicjatywy organów nadzorczych, a nie systemowej transparentności *ex ante*. Choć obowiązek udokumentowania oceny ryzyka zachowuje minimalny standard kontroli, brak jego ujawnienia w rejestrze ogranicza możliwość zewnętrznej weryfikacji – zarówno przez organy innych państw, jak i przez pozostałych interesariuszy. Zgodnie z motywem 9, dostęp do dokumentacji miałyby „właściwe organy krajowe”. Nie jest to rozwiązanie optymalne z punktu widzenia ochrony podmiotów prywatnych, w tym radców prawnych oraz ich klientów.

c) Zagrożenie dla wymiaru sprawiedliwości

Z perspektywy samorządu radców prawnych jako zawodu zaufania publicznego powołanego do profesjonalnego świadczenia pomocy prawnej, w tym polegającej na reprezentowaniu klientów przed sądami, szczególnie istotną i wysoce ryzykowną konsekwencją planowanej zmiany jest wyłączenie spośród obowiązku rejestracji – systemów AI stosowanych w wymiarze sprawiedliwości. Mając na uwadze otwarty i nieostry charakter wyjątków określonych w art. 6 ust. 3 AIA należy stwierdzić, że pozostawienie ich samodzielnej oceny w tym zakresie dostawcom stwarza ryzyko dla praw podstawowych, w tym prawa do sądu oraz dla fundamentalnych zasad praworządności oraz państwa prawa. Są to ryzyka szczególnie istotne z punktu widzenia prawnej ochrony interesów uczestników postępowań, na rzecz których radcowie prawni świadczą pomoc prawną. Dodać należy, że ewentualne szkody związane z wadliwym funkcjonowaniem systemów nie podlegających obowiązkowi rejestracji mogą być realnie trudne do naprawienia, biorąc pod uwagę nieodwracalne skutki niektórych decyzji procesowych, które mogą zostać podjęte w efekcie stosowania systemów. Pozostawienie wymogu rejestracji systemów nie uznanych za systemy wysokiego ryzyka pozwala na częściowe ograniczenie tych ryzyk *ex ante*.

2.2.4. Współpraca z organami ochrony praw podstawowych (art. 77 AIA)

Zgodnie z art. 77 AIA krajowe organy lub podmioty publiczne, które nadzorują lub egzekwują przestrzeganie obowiązków wynikających z prawa UE w zakresie ochrony praw podstawowych (tzw. FRABs – Fundamental Rights Authorities/Bodies) mają uprawnienie do bezpośredniego żądania dokumentacji od dostawców i podmiotów stosujących systemy AI wysokiego ryzyka, gdy dostęp do niej jest niezbędny do skutecznego wypełniania przez FRABs ich ustawowych zadań („mandatów”). Zgodnie z propozycją Digital Omnibus on AI – wskazane żądanie dostępu FRABs ma odbywać się poprzez pośrednika tj. organy nadzoru rynku (MSA).

Przykładowo w kontekście krajowym (zgodnie z projektem ustawy o systemach sztucznej inteligencji) głównym MSA ma być Komisja Rozwoju i Bezpieczeństwa Sztucznej Inteligencji (KRIBSI), choć w zależności od sektora mogą wchodzić w grę inne właściwe organy. Do FRABs w Polsce należą m.in. UODO, Rzecznik Praw Pacjenta, Państwowa Inspekcja Pracy i Rzecznik Praw Dziecka. Zgodnie z Rozporządzeniem 2024/1689 wykaz jest na bieżąco aktualizowany.

Kierunkowo zmianę tę należy ocenić pozytywnie, jednak z pewnymi zastrzeżeniami.

a) Pozytywne aspekty

Konstrukcja MSA jako scentralizowanego punktu kontaktowego może w założeniu przynieść liczne korzyści takie jak choćby lepsza koordynacja działań między organami i ustandaryzowana procedura (w tym jeden punkt kontaktowy zamiast wielu organów). Nowy obowiązek ścisłej współpracy i wzajemnej pomocy (proponowany art. 77 ust. 1b) może także usprawnić wymianę informacji (szczególnie w przypadkach transgranicznych).

Należy również pozytywnie ocenić propozycję zniesienia ograniczenia stosowania art. 77 ust. 1 AIA wyłącznie do systemów wysokiego ryzyka z Załącznika III, co naturalnie daje możliwość szerszej ochrony praw podstawowych.

b) Rola MSA w ocenie wniosków.

Proponowana regulacja nie określa w sposób dostatecznie precyzyjny, jaką rolę ma pełnić MSA przy przyjęciu wniosku. W naszej ocenie rola MSA powinna ograniczać się do

funkcji organizacyjno – koordynacyjnej (przyjmowania wniosku, zapewnienia kanału komunikacji z dostawcą i procedowania go przede wszystkim w aspektach wymogów formalnych – tworząc swego rodzaju “hub” administracyjny w tym obszarze). Merytoryczna ocena (w tym sama kwestia zasadności wniosku) powinna pozostać w kompetencji FRABs, co należy doprecyzować w propozycji art. 77 ust. 1 AIA.

Pozostawienie proponowanej konstrukcji tworzy ryzyko, że pośrednictwo MSA stanie się pewnego rodzaju mechanizmem filtrującym i blokującym, który może realnie osłabiać działanie FRABs i co za tym idzie efektywność ochrony praw podstawowych.

c) Ryzyko ograniczenia zakresu przekazywanych informacji

Propozycja art. 77 ust. 1a AIA zmienia także istotnie zakres uprawnień FRABs – z żądania dokumentacji „sporządzonej lub prowadzonej na mocy niniejszego rozporządzenia” (tj. przez operatora zgodnie z art. 77 ust. 1 AIA) na dokumentację „sporządzoną lub prowadzoną przez organ nadzoru rynku”. Przewiduje również, że MSA zwraca się do operatora o informacje „w razie potrzeby”, co *de facto* może pozostawiać MSA swobodną ocenę, czy zwrócenie się jest konieczne. Regulacja Digital Omnibus on AI nie daje też możliwości samodzielnego zwrócenia się do operatora, zatem istnieje ryzyko, że zakres informacji jaki otrzyma FRABs może być pośrednio kształtowany przez MSA.

d) Ryzyko nieefektywności i opóźnień

EDPB i EDPS we wspólnej opinii (1/2026) zauważają, że wymóg uzyskiwania informacji wyłącznie przez MSA może w praktyce prowadzić do nieefektywności i opóźnień w działaniach FRABs. Dodatkowy pośrednik w procedurze może wydłużyć czas reakcji na potencjalne naruszenia praw podstawowych. EDPB i EDPS postulują, by regulacja określała, że organy nadzoru rynku mają obowiązek przekazywać informacje żądane przez FRABs bez zbędnej zwłoki (zarówno na poziomie krajowym, jak i transgranicznym).

2.2.5. Propozycja wprowadzenia do AIA art. 4a tj. zmiany w zakresie przetwarzania danych wrażliwych

Zgodnie z art. 10 ust. 5 AIA dostawcy systemów AI wysokiego ryzyka mogą wyjątkowo przetwarzać szczególne kategorie danych osobowych w zakresie w jakim jest to bezwzględnie konieczne (*strictly necessary*) do zapewnienia wykrywania i korygowania stronniczości takich systemów – zgodnie z wymogami art. 10 ust. 2 lit. f) i g) AIA.

Warunkiem takiego przetwarzania jest stosowanie odpowiednich zabezpieczeń w zakresie podstawowych praw i wolności osób fizycznych.

Propozycja Digital Omnibus on AI zakłada wprowadzenie nowego przepisu tj. art. 4a, który miałby zastąpić przywołany wyżej art. 10 ust. 5 AIA, a także – co istotne – rozszerza znacząco jego zakres podmiotowy i przedmiotowy.

a) Zakres podmiotowy i przedmiotowy

Projekt rozszerza mechanizm dopuszczający przetwarzanie szczególnych kategorii danych osobowych na potrzeby wykrywania i korygowania stronniczości nie tylko na dostawców, lecz również na podmioty stosujące. Jednocześnie, zgodnie z art. 4a ust. 2, rozwiązanie to znajdzie zastosowanie także w odniesieniu do szerszego katalogu systemów i modeli AI (nie tylko systemów wysokiego ryzyka, lecz również innych systemów oraz modeli, w tym modeli GPAI), co istotnie zwiększa liczbę podmiotów i „kontekstów”, w których dopuszczalne byłoby przetwarzanie danych wrażliwych. Rozszerzenie kręgu podmiotów budzi wątpliwości, ponieważ podmioty stosujące – w odróżnieniu od dostawców – w wielu przypadkach nie mają pełnej wiedzy o architekturze systemu AI ani dostępu do informacji o danych wykorzystanych w procesie uczenia, co może utrudniać rzetelną ocenę, czy przetwarzanie danych wrażliwych jest rzeczywiście niezbędne do wykrycia i skorygowania stronniczości.

b) Standard konieczności i stanowisko EDPB i EDPS

Obecny art. 10 ust. 5 AIA wymaga, aby przetwarzanie danych wrażliwych było „ściśle niezbędne” (strictly necessary). Proponowana przez Digital Omnibus on AI zmiana obniża ten standard tj. art. 4a ust. 1 posługuje się kryterium „niezbędne” (necessary), natomiast art. 4a ust. 2 kryterium „niezbędne i proporcjonalne” (necessary and proportionate). Złagodzenie wymogu niezbędności budzi wątpliwości w świetle dość restrykcyjnego podejścia prawa UE do przetwarzania szczególnych kategorii danych osobowych, które jest co do zasady zakazane na gruncie art. 9 ust. 1 RODO, a wyjątki od tego zakazu powinny być definiowane wąsko. EDPB i EDPS we wspólnej opinii 1/2026 zwracają uwagę, że o ile cel walki ze stronniczością jest słuszny, o tyle propozycja Komisji osłabia standard ochrony, zmieniając wymóg „ściślej niezbędności” na zwykłą „niezbędność”.

Organy wskazują również, że aby ograniczyć ryzyko nadużyć, przypadki, w których dostawcy i podmioty stosujące – mogłyby powoływać się na wyjątek przewidziany w art. 4a w odniesieniu do systemów i modeli AI innych niż wysokiego ryzyka – powinny być

jasno określone i zawężone do sytuacji, w których ryzyko istotnych negatywnych skutków wynikających ze stronniczości jest rzeczywiście poważne.

c) Normalizacja masowego gromadzenia danych

Wprowadzenie art. 4a w proponowanym brzmieniu może prowadzić do normalizacji przetwarzania (a w praktyce także pozyskiwania i agregowania) szczególnych kategorii danych osobowych pod szeroko ujętym celem „wykrywania i korygowania stronniczości”. Rozszerzenie tego wyjątku na wszystkie systemy i modele AI (nie tylko wysokiego ryzyka) oraz na szeroki krąg podmiotów (w tym podmioty stosujące) zwiększa ryzyko, że przetwarzanie danych wrażliwych stanie się w praktyce rozwiązaniem „domyślnym” lub „na zapas” (zwłaszcza że w pierwszej kolejności ocena, czy przetwarzanie danych wrażliwych jest rzeczywiście niezbędne, będzie w dużej mierze pozostawiona wielu różnym podmiotom, co utrudnia zapewnienie jednolitych standardów i skuteczną weryfikację zasadności takich działań).

W tym kontekście Irlandzka Komisja Praw Człowieka i Równości (IHREC) wskazuje na ryzyko, że nowy art. 4a poszerza dostęp do danych wrażliwych w sposób, który może sprzyjać ich masowemu gromadzeniu, a w konsekwencji ułatwiać profilowanie i praktyki o charakterze nadzorczym pod pretekstem „bias detection”, przy niedostatecznie silnych mechanizmach egzekwowania. W rezultacie mechanizm, który miał służyć ograniczaniu dyskryminacji, może paradoksalnie zwiększyć skalę operacji na danych szczególnie wrażliwych oraz utrudnić skuteczną kontrolę nad ich zakresem i celowością (zob. Letter to Department of Enterprise, Tourism and Employment on EU Interpretative Note Art 77 bodies and Digital Omnibus, 8.12.2025).

d) Ocena krytyczna

Z perspektywy ochrony praw podstawowych proponowane zmiany budzą uzasadnione obawy. Dane wrażliwe (pochodzenie rasowe lub etniczne, poglądy polityczne, przekonania religijne, dane dotyczące zdrowia, orientacji seksualnej itp.) podlegają szczególnej ochronie właśnie ze względu na potencjał dyskryminacyjny ich przetwarzania. Rozszerzenie możliwości przetwarzania tych danych na szerszy katalog systemów AI i podmiotów, przy jednoczesnym obniżeniu standardu ścisłej niezbędności, stwarza ryzyko, że wyjątek przewidziany na potrzeby walki z dyskryminacją może zostać wykorzystany w sposób sprzeczny z tym celem.

V. Podsumowanie i wnioski

Pakiet Digital Omnibus i Digital Omnibus on AI, mimo deklarowanego celu uproszczenia i zwiększenia konkurencyjności europejskiej, zawiera propozycje budzące poważne wątpliwości z perspektywy ochrony praw podstawowych i spójności systemu prawnego UE. Szczególnie niepokojące są zmiany w RODO, które pod pozorem „technicznego doprecyzowania” mogą prowadzić do istotnego osłabienia ochrony danych osobowych.

W ocenie Ekspertów Komisji Nowych Technologii Krajowej Rady Radców Prawnych należy rozważyć:

- 1) Odstąpienie od zmian definicji danych osobowych w art. 4 pkt 1 RODO, które subiektywizują obiektywne dotychczas pojęcie i zniekształcają logikę motywu 26.
- 2) Rezygnację z wprowadzania art. 88c RODO jako naruszającego zasadę neutralności technologicznej, sprzecznego z orzecznictwem TSUE i zbędnego w świetle istniejącej wykładni EDPB (opinia 28/2024).
- 3) Rezygnację z proponowanego art. 9 ust. 2 lit. k) RODO i utrzymanie regulacji przetwarzania szczególnych kategorii danych dla celów wykrywania uprzedzeń w art. 10 ust. 5 AIA z zachowaniem standardu „ściślej konieczności”.
- 4) Zachowanie zakazowego charakteru art. 22 RODO oraz wymogu „niezbędności” w art. 22 ust. 2 lit. a) RODO jako fundamentalnej gwarancji ochrony przed zautomatyzowanym podejmowaniem decyzji.
- 5) Odstąpienie od rozszerzania wyłączeń obowiązków informacyjnych w art. 13 RODO ponad istniejące już wyłączenie w art. 13 ust. 4.
- 6) Skoncentrowanie działań na wzmocnieniu egzekwowania istniejących przepisów RODO oraz wydawaniu wytycznych przez EDPB, zamiast nowelizacji materialnych przepisów rozporządzenia.
- 7) Przeprowadzenie pełnej oceny wpływu na prawa podstawowe przed przyjęciem proponowanych zmian – zarówno w zakresie Digital Omnibus, jak i Digital Omnibus on AI.

- 8) W zakresie proponowanych zmian co do art. 111 i 113 AIA (harmonogram stosowania przepisów) – zastąpienie mechanizmu ruchomego terminu (art. 113 AIA) – sztywną datą wejścia w życie przepisów w celu usunięcia niepewności prawnej. Ewentualnie, jeśli mechanizm ruchomy miałby zostać utrzymany – określenie terminu, w którym Komisja zobowiązana jest wydać decyzję, tak aby podmioty zobowiązane dysponowały minimalnym gwarantowanym okresem na wdrożenie zmian. Należy rozważyć także zróżnicowanie odroczenia według kategorii obowiązków tj. wymogi, których stosowanie nie jest uzależnione od dostępności szczegółowych standardów technicznych (np. obowiązki w zakresie przejrzystości i informowania), mogłyby wejść w życie zgodnie z pierwotnym harmonogramem.
- 9) Utrzymanie obowiązkowego charakteru przepisów dotyczących AI literacy (Art. 4 AIA), zgodnie z rekomendacją EDPB i EDPS. Zamiast znoszenia obowiązku, należy rozważyć wydanie przez Komisję praktycznych wytycznych dla dostawców i podmiotów stosujących AI w zakresie wdrażania tego wymogu.
- 10) Wprowadzenie w projektowanym art. 77 AIA ust. 1 a wymogu działania MSA „bez zbędnej zwłoki” (zarówno w przypadkach na poziomie krajowym, jak i transgranicznym) oraz doprecyzowanie, że merytoryczna ocena zasadności wniosku powinna pozostać w kompetencji FRABs.
- 11) Utrzymanie obowiązku rejestracji w unijnej bazie danych systemów AI, co do których dostawca uznał, że nie są systemami wysokiego ryzyka.
- 12) Pozostawienie obecnego art. 10 ust. 5 AIA i rezygnację z art. 4a jako rozwiązania rozszerzającego przetwarzanie danych wrażliwych. Alternatywnie, jeżeli art. 4a miałby zostać utrzymany, należy rozważyć pozostawienie go jako ograniczonego do systemów wysokiego ryzyka z przywróceniem wymogu „ściśle niezbędności” (strictly necessary), a jego stosowanie poza kategorią systemów wysokiego ryzyka powinno być dopuszczalne wyłącznie w ściśle określonych przypadkach poważnego ryzyka stronnictwośi, przy wzmocnionych wymogach dokumentacyjnych i nadzorczych.
- 13) W toku dalszych prac legislacyjnych nad wnioskiem – jako zasadne wydaje się przeprowadzenie szerszych konsultacji z udziałem zrównoważonej reprezentacji wszystkich zainteresowanych stron, w tym organizacji społeczeństwa

obywatelskiego, środowisk naukowych, organów ochrony praw podstawowych oraz samorządów zawodowych.

Prawdziwe uproszczenie nie polega na przepisywaniu praw, lecz na jasnym egzekwowaniu istniejących regulacji i silniejszym nadzorze. Wiarygodność Europy jako obrońcy praw cyfrowych zależy od utrzymania ochrony, którą zbudowała, a nie od jej osłabiania pod presją deregulacyjną.

Wyrażamy nadzieję, że powyższe uwagi zostaną wzięte pod uwagę w dalszym procesie legislacyjnym.



KRAJOWA IZBA
RADCÓW PRAWNYCH

Warszawa, luty 2026