

RAPORT

AI W ADMINISTRACJI PUBLICZNEJ

**Priorytety Transformacji Cyfrowej w Polsce
Wyzwania i rekomendacje na lata 2026-2029
Perspektywa Obywatela**



krajowa
izba
gospodarki
cyfrowej

AI w administracji publicznej

Priorytety Transformacji Cyfrowej w Polsce
Wyzwania i rekomendacje na lata 2026-2029
Perspektywa Obywatela

OPIEKA MERYTORYCZNA

Iwona Wendel

DTP I PRACE GRAFICZNE:

DigitalH Sp. z o.o.

AUTORZY RAPORTU:

dr inż. Karol Antczak
Piotr Bieszczad
Maciej Cieśla
Krzysztof Malik
Adam Paczusi

REDAKCJA I KOREKTA

Katarzyna Kuczyńska



RAPORT KIGC 2026

AI w administracji publicznej

Priorytety Transformacji Cyfrowej w Polsce
Wyzwania i rekomendacje na lata 2026-2029
Perspektywa Obywatela

PARTNER MERYTORYCZNY



PATRON MEDIALNY

CyberDefence **24**

PATRON TECHNOLOGICZNY

DELL
Technologies

AI w administracji publicznej

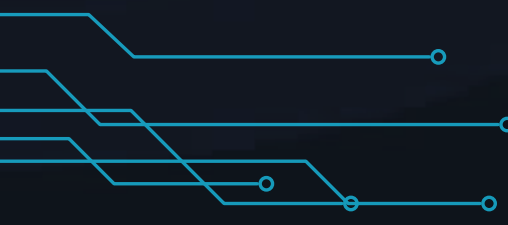
Szanowni Państwo,

transformacja cyfrowa państwa wchodzi dziś w kolejny, przełomowy etap. Sztuczna inteligencja, która jeszcze niedawno była postrzegana jako technologia przyszłości, staje się istotnym elementem codziennego funkcjonowania społeczeństwa oraz coraz wyraźniej zaznacza swoją rolę w działaniach administracji publicznej.

Z perspektywy Ministerstwa Cyfryzacji widzimy, że zmiana ta zachodzi równolegle na wielu poziomach. Dynamiczny rozwój technologii stwarza realną możliwość automatyzacji procesów, lepszego wykorzystania danych oraz budowy nowoczesnych usług publicznych. Jednocześnie rosną oczekiwania obywateli wobec państwa, nie tylko w zakresie dostępności usług cyfrowych, ale także ich jakości, szybkości oraz przejrzystości. Transformacja ta ma również istotne znaczenie dla budowania odporności państwa oraz jego zdolności do reagowania na współczesne wyzwania. Mamy jednocześnie pełną świadomość, że wdrażanie sztucznej inteligencji w sektorze publicznym wiąże się ze szczególną odpowiedzialnością. Państwo przetwarza dane o najwyższym poziomie wrażliwości, a ich ochrona pozostaje jednym z fundamentów zaufania społecznego. W warunkach rosnącej liczby cyberataków, pojawiania się nowych form zagrożeń, takich jak deep-fake czy dezinformacja, oraz zwiększającej się złożoności systemów cyfrowych, bezpieczeństwo musi pozostać absolutnym priorytetem.

Dlatego podejście do rozwoju sztucznej inteligencji w administracji publicznej wymaga równowagi. Z jednej strony konieczne jest wykorzystanie potencjału tej technologii do zwiększenia efektywności działania państwa i poprawy jakości usług publicznych. Z drugiej strony wdrożenia te muszą odbywać się w sposób kontrolowany, zgodny z regulacjami oraz zapewniający pełną ochronę danych obywateli.

„Państwo przetwarza dane o najwyższym poziomie wrażliwości, a ich ochrona pozostaje jednym z fundamentów zaufania społecznego.”



PRZEDMOWA

Rolą Ministerstwa Cyfryzacji jest tworzenie warunków, w których rozwój nowych technologii może następować w sposób bezpieczny, spójny i odpowiedzialny. Obejmuje to zarówno rozwój usług i infrastruktury cyfrowej państwa, jak i budowę ram regulacyjnych oraz standardów, które umożliwiają instytucjom publicznym świadome korzystanie z nowych rozwiązań. W tym procesie kluczowe znaczenie ma dialog oraz współpraca. Rozwój sztucznej inteligencji w administracji publicznej wymaga uwzględnienia perspektywy obywateli, doświadczeń instytucji publicznych oraz wiedzy i kompetencji środowiska eksperckiego i rynku. Tylko takie podejście pozwala budować rozwiązania, które są jednocześnie innowacyjne, bezpieczne i dopasowane do realnych potrzeb.

Jesteśmy przekonani, że odpowiedzialne i przemyślane wykorzystanie sztucznej inteligencji może stać się jednym z kluczowych elementów dalszego rozwoju administracji publicznej w Polsce. Warunkiem jest jednak podejście systemowe, oparte na bezpieczeństwie, zaufaniu oraz współpracy.



Paweł Olszewski

Sekretarz Stanu

Ministerstwo Cyfryzacji

Executive Summary: AI w administracji publicznej – jak wdrażać bezpiecznie i skutecznie

Poniżej prezentujemy kluczowe wnioski oraz rekomendacje płynące z niniejszego raportu, które mają wesprzeć administrację publiczną i jednostki samorządowe w odpowiedzialnym, bezpiecznym i efektywnym wdrażaniu rozwiązań opartych na sztucznej inteligencji.

AI już działa – ale poza kontrolą systemową

Dane jednoznacznie pokazują, że transformacja już się rozpoczęła – i to oddolnie:

69 %

Polaków korzysta z AI regularnie

54 %

pracowników używa AI
w pracy zawodowej

46 %

urzędników wykorzystuje gen. AI
w codziennej pracy

90 %

nawet tyle użytkowników korzysta z publicznych
narzędzi AI (często darmowych)

Jednocześnie:

70 %

urzędników wskazuje brak procedur
jako główną barierę

55 %

pracowników nie ujawnia korzystania
z AI w pracy

W praktyce oznacza to możliwość wdrożenia rozwiązań, które bezpośrednio przekładają się na jakość działania instytucji publicznych. Sztuczna inteligencja może wspierać m.in. analizę i streszczanie dokumentów, automatyczne przygotowywanie odpowiedzi, szybkie wyszukiwanie informacji w aktach czy obsługę obywateli w różnych językach. Coraz większe znaczenie ma również wykorzystanie AI w zapleczu technologicznym administracji – w tym w utrzymaniu systemów IT czy tworzeniu oprogramowania.

AI pozwala nie tylko przyspieszyć procesy, ale także zwiększyć ich dostępność i spójność. Skrócenie czasu obsługi spraw, ograniczenie błędów oraz lepszy dostęp do informacji przekładają się bezpośrednio na doświadczenie obywatela i efektywność działania urzędu.

Rekomendacja:

Kluczowym zadaniem administracji jest nie tyle ograniczanie dostępu do technologii, ile stworzenie bezpiecznych, kontrolowanych alternatyw. Oznacza to wdrażanie własnych rozwiązań opartych na sztucznej inteligencji – w szczególności modeli rozwijanych lokalnie i działających w krajowej infrastrukturze – które zapewnią pełną kontrolę nad danymi, zgodność z regulacjami oraz realne wsparcie pracy urzędników.

Korzyści z wdrażania sztucznej inteligencji w administracji publicznej

Sztuczna inteligencja nie zastępuje urzędników – wzmacnia ich efektywność i zmienia sposób wykonywania codziennych obowiązków. Jej największą wartością nie jest automatyzacja całych stanowisk pracy, lecz przejmowanie powtarzalnych, czasochłonnych czynności, które dziś obciążają administrację i wydłużają obsługę obywateli.

Dane pokazują jednoznacznie, że pracownicy administracji dostrzegają w AI przede wszystkim narzędzie wspierające, a nie zastępujące ich rolę:

81 %

urzędników widzi potencjał
w automatyzacji zadań

76 %

wskazuje na realną oszczędność czasu

72 %

dostrzega usprawnienie pracy
z dokumentami

77 %

podkreśla znaczenie AI
w tłumaczeniach i komunikacji



Raport KIGC 2026

W praktyce oznacza to możliwość wdrożenia rozwiązań, które bezpośrednio przekładają się na jakość działania instytucji publicznych. Sztuczna inteligencja może wspierać m.in. analizę i streszczanie dokumentów, automatyczne przygotowywanie odpowiedzi, szybkie wyszukiwanie informacji w aktach czy obsługę obywateli w różnych językach. Coraz większe znaczenie ma również wykorzystanie AI w zapleczu technologicznym administracji – w tym w utrzymaniu systemów IT czy tworzeniu oprogramowania.

W efekcie AI pozwala nie tylko przyspieszyć procesy, ale także zwiększyć ich dostępność i spójność. Skrócenie czasu obsługi spraw, ograniczenie błędów oraz lepszy dostęp do informacji przekładają się bezpośrednio na doświadczenie obywatela i efektywność działania urzędu.

Rekomendacja:

Administracja publiczna powinna koncentrować wdrożenia AI na konkretnych, wysokowartościowych obszarach operacyjnych – w szczególności tam, gdzie technologia może szybko odciążyć pracowników i poprawić jakość obsługi obywatela. Priorytetem powinny być rozwiązania wspierające pracę z dokumentami, komunikację oraz dostęp do informacji, wdrażane w sposób kontrolowany i zintegrowany z istniejącymi systemami.

Ryzyka obecnego stanu rzeczy: bezpieczeństwo, brak kontroli i odpowiedzialność osobista pracowników

Największe ryzyko związane z wdrażaniem sztucznej inteligencji w administracji publicznej nie wynika z samej technologii, lecz z braku kontroli nad jej wykorzystaniem oraz rosnącej odpowiedzialności – nie tylko instytucjonalnej, ale również osobistej pracowników i kadry zarządzającej.

Wraz z rosnącą skalą wykorzystania AI zwiększa się również powierzchnia zagrożeń. Administracja publiczna już dziś należy do najbardziej atakowanych sektorów, a rozwój AI dodatkowo wzmacnia aktywność cyberprzestępców. Rosnąca liczba ataków ransomware, wycieków danych czy nowych form zagrożeń – takich jak deepfake i AI-asystowana dezinformacja – sprawia, że bezpieczeństwo cyfrowe staje się jednym z kluczowych wyzwań funkcjonowania państwa.



Raport KIGC 2026

Jednocześnie zmienia się charakter odpowiedzialności za te obszary. Nowe regulacje, w tym AI Act oraz dyrektywa NIS 2, wyraźnie przesuwają ciężar odpowiedzialności na poziom organizacyjny i osobowy. Oznacza to, że konsekwencje błędów – takich jak niewłaściwe wykorzystanie narzędzi AI, brak nadzoru nad systemami czy niedostateczne zabezpieczenie danych – mogą dotyczyć nie tylko instytucji, ale również konkretnych osób zarządzających i pracowników. W praktyce oznacza to ryzyko wysokich kar finansowych oraz odpowiedzialności za niedochowanie należytej staranności.

Na poziomie operacyjnym szczególnie istotne pozostają:

- brak kontroli nad przetwarzaniem informacji i ich lokalizacją,
- podejmowanie lub wspieranie decyzji administracyjnych przez systemy AI bez odpowiedniego nadzoru człowieka,
- niezgodność działań z regulacjami (AI Act, RODO, NIS 2).

W obecnym stanie rzeczy pracownicy często działają w warunkach niepewności – bez jasnych procedur i narzędzi, które pozwalałyby korzystać z AI w sposób bezpieczny i zgodny z prawem. To powoduje, że indywidualne decyzje operacyjne mogą nieświadomie generować istotne ryzyka dla całej organizacji.

Rekomendacja:

Administracja publiczna powinna pilnie wdrożyć jasne ramy zarządzania wykorzystaniem AI, które nie tylko zwiększą poziom bezpieczeństwa, ale również ochronią pracowników przed nieświadomym naruszeniem przepisów. Kluczowe jest zapewnienie bezpiecznych narzędzi, precyzyjnych procedur oraz systemowych szkoleń, tak aby odpowiedzialność za wykorzystanie AI była realnie wspierana przez organizację, a nie przerzucana na pojedynczych urzędników.

Rola ekosystemu i partnerstw we wdrażaniu AI

Skala i złożoność wdrażania sztucznej inteligencji w administracji publicznej sprawiają, że proces ten nie powinien być realizowany wyłącznie siłami pojedynczych instytucji. Kluczowe znaczenie ma wykorzystanie doświadczenia rynku oraz współpraca z wyspecjalizowanymi podmiotami, które łączą kompetencje technologiczne, prawne i operacyjne.



W praktyce oznacza to budowanie modelu wdrożeń opartego na zaufanych partnerach – organizacjach branżowych, ekspertach oraz firmach posiadających doświadczenie w implementacji rozwiązań AI w środowiskach wymagających wysokiego poziomu bezpieczeństwa i zgodności regulacyjnej.

Takie podejście pozwala:

- skrócić czas wdrożeń i ograniczyć ryzyko błędów,
- korzystać z już sprawdzonych standardów i dobrych praktyk,
- zapewnić zgodność z dynamicznie zmieniającymi się regulacjami,
- budować rozwiązania dopasowane do realiów administracji publicznej.

Rekomendacja:

Administracja publiczna powinna aktywnie wykorzystywać wsparcie ekspertów rynkowych i organizacji branżowych przy projektowaniu i wdrażaniu rozwiązań AI. Budowa zaufanego ekosystemu partnerów – łączącego kompetencje technologiczne, regulacyjne i operacyjne – jest jednym z kluczowych warunków bezpiecznej i skutecznej transformacji cyfrowej państwa.

1. Obywatel w świecie nowych technologii

- Obywatel w centrum transformacji cyfrowej 19
- Obywatel w świecie nowych technologii 21
- AI w codziennym życiu obywateli 22
- Oczekiwania obywateli wobec administracji 23
- Wnioski: obywatel jako punkt odniesienia 25

2. Kontekst technologiczny

- AI i transformacja pracy 29
- Rozwój AI i modeli językowych 33

3. Regulacje i zmiany w prawie

- Regulacje i odpowiedzialność - wprowadzenie 41
- Akt w sprawie sztucznej inteligencji (AI Act) 42
- Dyrektywa NIS 2 47
- RODO w kontekście AI 53
- KSC i krajowe regulacje 58
- Prawo Zamówień Publicznych (PZP) 61

4. AI w administracji publicznej

- AI w administracji publicznej - przegląd 67
- Obsługa obywateli 71
- Wsparcie pracy urzędników oraz działów IT 73
- Wnioski operacyjne 77

5. Bezpieczeństwo danych i AI

- Bezpieczeństwo danych i AI - wprowadzenie 80
- Dane wrażliwe w administracji 83
- Przetwarzanie dokumentów 84
- Ryzyka publicznych modeli AI 85
- Kontrola infrastruktury 86
- Audyt i zgodność 89
- Scenariusze ryzyk 91
- Wnioski bezpieczeństwa 93

6. Studium przypadków

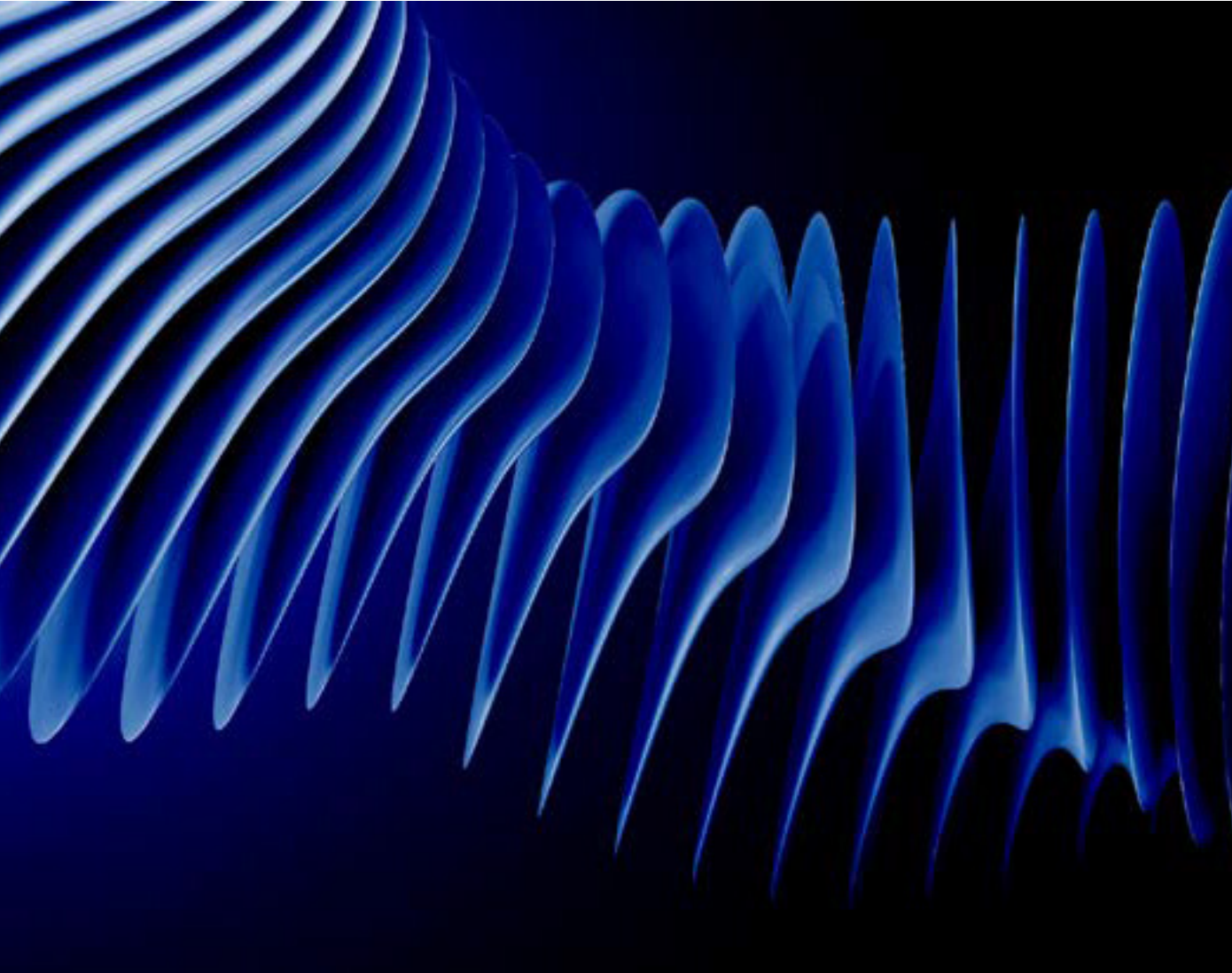
- Case Studies - Safe VOX 96
- Case Studies - Safe TAIA 98
- Case Studies - Safe CODER 102

7. Suwerenność technologiczna

• Suwerenność technologiczna - wprowadzenie	106
• Suwerenna chmura	108
• Lokalne modele AI	111
• Infrastruktura państwowa	116
• Zależność od globalnych dostawców i ryzyko outsourcingu AI	118
• Rola krajowych technologii	120
• Wnioski strategiczne	122

8. Rekomendacje

• Rekomendacje i wnioski z raportu - wersja rozszerzona	126
---	-----



Wstęp: Sztuczna inteligencja i państwo nowej ery

Szanowni Czytelnicy, oddajemy w Państwa ręce raport „AI w administracji publicznej. Raport 2026”. Żyjemy w czasach bezprecedensowego przyspieszenia technologicznego. Sztuczna inteligencja (AI) przestała być domeną akademickich laboratoriów badawczych i wizji z literatury science fiction, stając się fundamentem funkcjonowania nowoczesnego państwa. Skala tego zjawiska jest fascynująca: aż 69% Polaków korzysta z AI w sposób regularny⁰¹. Co więcej, cyfrowa rewolucja na dobre dotarła do serca polskiego aparatu państwowego – już 46% pracowników administracji publicznej wykorzystuje na co dzień narzędzia generatywnej sztucznej inteligencji⁰².

Jako państwo i społeczeństwo стоимy jednak przed potężnym, wielowymiarowym wyzwaniem. Administracja publiczna musi sprostać rosnącym oczekiwaniom obywateli, jednocześnie

mierząc się z problemem regulacji, bezpieczeństwa danych, zmian strukturalnych i demograficznych. Sztuczna inteligencja daje niepowtarzalną szansę na automatyzację rutynowych

zadań i ułatwienie nawigowania w gąszczu zawitych przepisów. Z drugiej strony, jej niekontrolowane wdrożenie generuje olbrzymie zagrożenia. Urzędy państwowe przetwarzają najbardziej wrażliwe dane o obywatelach. Tymczasem badania rynkowe ujawniają alarmujące zjawisko tzw. „shadow AI” – w celu optymalizacji pracy wykorzystuje się publicznie dostępne, darmowe narzędzia, często nieświadomie ryzykując wyciekiem poufnych informacji. Wyzwaniem dla Polski nie jest więc już odpowiedź na pytanie, „czy” wdrażać AI w urzędach, lecz „jak” dokonać tego, by drastycznie zwiększyć ich efektywność, nie idąc na najmniejszy kompromis w kwestii bezpieczeństwa danych i utrzymania suwerenności technologicznej. Uzależnienie infrastruktury państwa od globalnych gigantów cyfrowych (Big Tech) może być bowiem równie niebezpieczne, co historyczne uzależnienie od importu zagranicznych surowców energetycznych.

Raport obejmuje perspektywę obywateli, zastosowania AI, regulacje, bezpieczeństwo i suwerenność technologiczną. Zaczynamy od postawienia w centrum samego człowieka. W sekcji „Obywatel w świecie nowych technologii” przyglądamy się, jak cyfryzacja zmienia oczekiwania względem relacji na linii państwo-mieszkaniec. Pokolenie wchodzące na rynek pracy nie wyobraża sobie funkcjonowania bez wsparcia AI, choć transformacja ta budzi też lęki – 16% badanych urzędników obawia się, że algoryt-

my mogą z czasem zastąpić ich na stanowiskach pracy⁰³. Zrozumienie tych barier jest kluczem do udanej zmiany. Następnie, w rozdziale poświęconym „Sztucznej Inteligencji w administracji publicznej”, analizujemy technologiczne podstawy wykorzystania AI w administracji publicznej. Pokazujemy, jak państwo zamierza odciążyć urzędników dzięki inteligentnej analizie dokumentów i zautomatyzowanej obsłudze. Opieramy się tu na polskich, bezpiecznych i otwartych modelach językowych, takich jak wielofunkcyjny PLLuM, który wkrótce zasili aplikację mObywatel w formie wirtualnego asystenta.

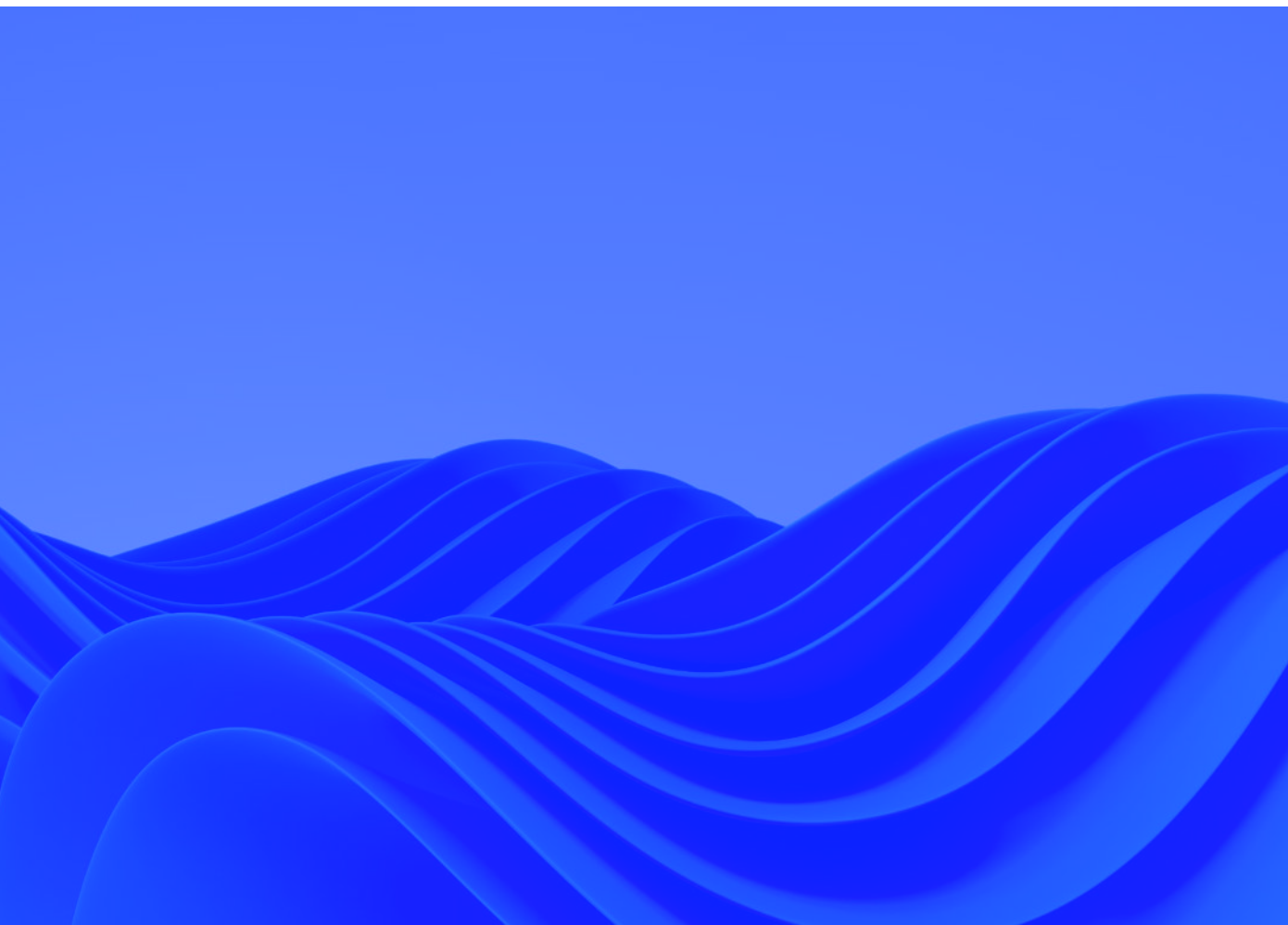
Technologia to jednak tylko narzędzie – aby służyła społeczeństwu, musi funkcjonować w jasno określonych ramach prawnych. W części „Bezpieczeństwo danych i zgodność regulacyjna” rozkładane są na czynniki pierwsze legislacyjne filary cyfrowego państwa. Wyjaśniamy, w jaki sposób europejski akt w sprawie sztucznej inteligencji (AI Act) klasyfikuje systemy używane m.in. w opiece zdrowotnej czy wymiarze sprawiedliwości jako rozwiązania „wysokiego ryzyka”, wymuszając rygorystyczny nadzór człowieka. Omawiamy też nową architekturę cyberbezpieczeństwa ukształtowaną przez dyrektywę NIS 2 i ustawę KSC oraz pokazujemy, jak unijne RODO dyktuje warunki bezpiecznego trenowania modeli AI poprzez mechanizmy minimalizacji danych. Tuż po tym przechodzimy do tematu „Suwerenności technologicznej”. Uzasadniamy,

dlaczego Polska strategicznie inwestuje we własne moce obliczeniowe, powołując do życia Fabryki AI (m.in. w Krakowie i Poznaniu) oraz współtworząc bałtycką Gigafabrykę AI. Rozwój krajowej infrastruktury chmurowej może zwiększać kontrolę nad przetwarzaniem danych i ograniczać zależność od zewnętrznych dostawców.

Aby jednak ta wiedza nie pozostała jedynie w sferze teorii, w rozdziale „Studia przypadków sektorowych” przenosimy dyskusję na grunt praktyki. Udowadniamy w ten sposób, jak algorytmy realnie wspierają medyków, skracają przewlekłość postępowań i automatyzują biurokrację. Całość spięta jest w „Rekomendacjach

dla urzędów, gmin i samorządów”, dostarczając decydom kompletne przewodniki, który krok po kroku instruuje, jak budować bezpieczną architekturę AI, zarządzać ryzykiem algorytmicznym i inwestować w cyfrowe kompetencje zespołów.

Niniejszy raport to merytoryczne i dowodowe kompendium, z którego płynie jednoznaczny wniosek: sztuczna inteligencja w administracji to nie zagrożenie dla człowieka, lecz historyczny skok ewolucyjny. To szansa na sprawne państwo, które dba o najwyższe standardy obsługi obywatela, szanując jego prywatność. Serdecznie zapraszamy do lektury!



Wstęp

1. Raport KPMG, AI Trust w Polsce 2025, KPMG, 2025.
2. Artykuł NASK – Państwowy Instytut Badawczy, Prawie połowa polskich urzędników korzysta z generatywnej AI, 2024.
3. PIE cytowany przez PAP (Polska Agencja Prasowa), na podstawie danych Polskiego Instytutu Ekonomicznego, 2024.

01

Obywatel w centrum transformacji cyfrowej



krajowa
izba
gospodarki
cyfrowej

Obywatel w centrum transformacji cyfrowej.

Krajobraz społeczny w trakcie algorytmicznej rewolucji

Obywatel to nie tylko ostateczny beneficjent usług publicznych, ale przede wszystkim nadrzędny punkt odniesienia dla wszelkich innowacji technologicznych wdrażanych przez aparat państwowy. Głębokie zrozumienie postaw, obaw, rzeczywistych kompetencji oraz precyzyjnie zdefiniowanych oczekiwań społeczeństwa stanowi absolutny fundament udanej i bezpiecznej transformacji cyfrowej. Bez szerokiej akceptacji społecznej i zaufania do państwa, nawet najbardziej zaawansowane, wielomiliardowe systemy algorytmiczne pozostaną jedynie bezużytecznymi, technokratycznymi projektami. Niniejszy rozdział przedstawia najważniejsze społeczne aspekty cyfryzacji i wykorzystania AI, zwiastując kluczowe tematy, które zostaną szczegółowo rozwinięte w kolejnych sekcjach raportu.

Od petenta do klienta cyfrowego

W kolejnym podrozdziale zatytułowanym „Obywatel w świecie nowych technologii” przeanalizowany zostanie obecny stan cyfryzacji polskiego społeczeństwa, które w ostatnich latach przeszło bezprecedensowy proces „wymuszonej dojrzałości cyfrowej”. Zmiana relacji na linii państwo-obywatel doprowadziła do trwałej ewolucji dawnego „petenta” w świadomego „klienta cyfrowego”. Jak udowadniają twarde dane statystyczne, zdecydowana większość polskiego społeczeństwa wysoce docenia cyfrowe usługi publiczne, uznając je za kolosalne ułatwienie w codziennym życiu. Jednocześnie jednak, ta technologiczna ewolucja ujawnia głębokie asymetrie kompetencyjne i pęknięcia międzypokoleniowe, które mogą prowadzić do cyfrowego wykluczenia części społeczeństwa, szczególnie osób starszych, na co administracja musi znaleźć szybką odpowiedź.

Zrozumienie tych zjawisk płynnie prowadzi do kolejnej sekcji: „AI w codziennym życiu obywateli”. Ukaże ona fascynujący, choć pełen sprzeczności obraz polskiego społeczeństwa w zderzeniu z rewolucją algorytmiczną. Z jednej strony, obywatele wykazują ogromny entuzjazm – niemal siedmiu na dziesięciu Polaków regularnie korzysta ze sztucznej inteligencji, wyprzedzając pod tym względem

wiele państw zachodnich. Z drugiej strony, analitycy określają ten stan mianem „krajobrazu pełnego paradoksów”. Masowej adaptacji nowinek technologicznych, najczęściej w postaci darmowych, publicznych platform, towarzyszy alarmująco niski poziom świadomości na temat zagrożeń dla prywatności, ryzyka wycieku danych czy wszechobecnej dezinformacji. W tej części raportu obnażone zostaną lęki obywateli przed nową technologią oraz systemowe braki w edukacji cyfrowej.

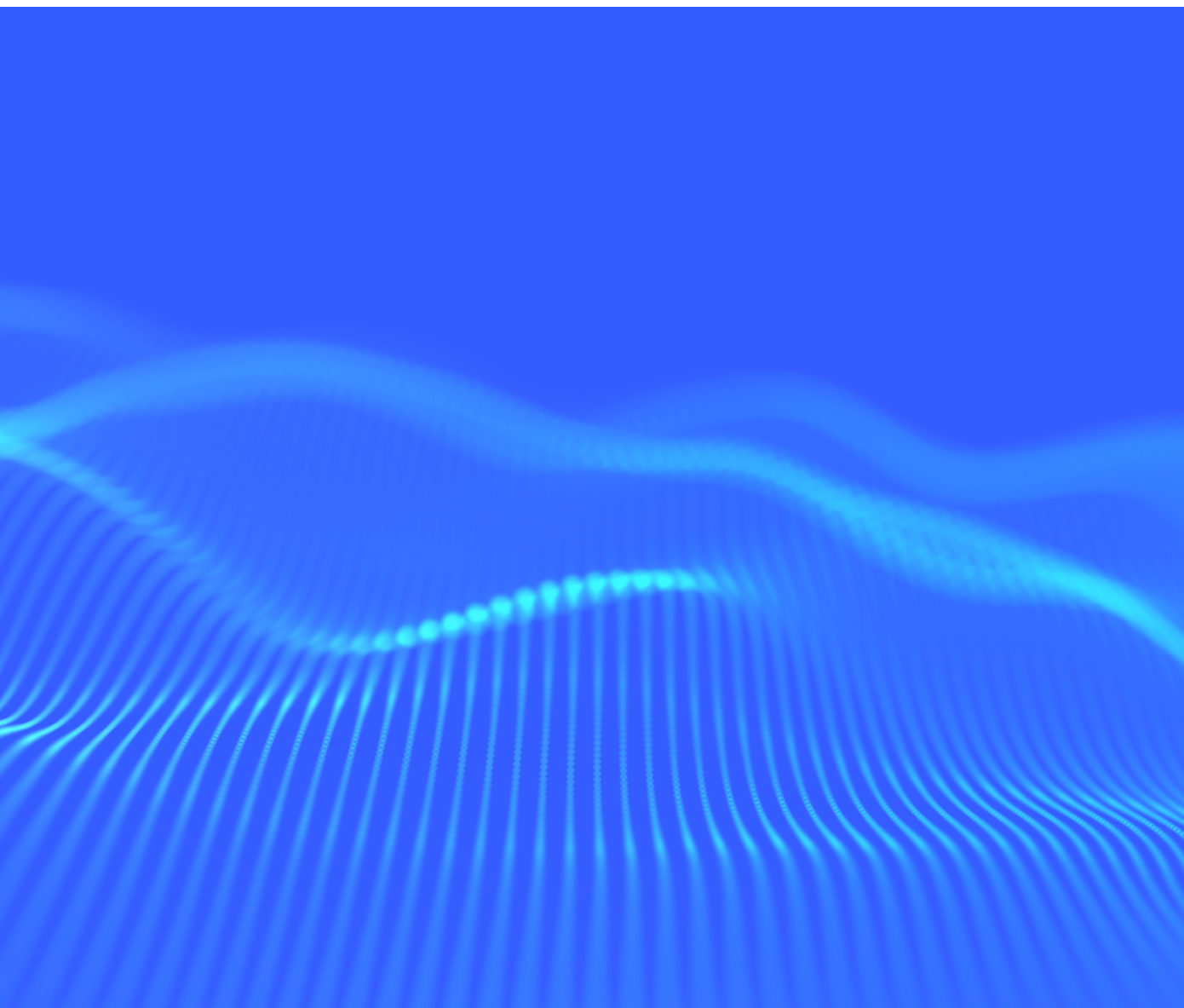
Sprecyzowane oczekiwania i wyznaczone granice

Te codzienne, cyfrowe doświadczenia kształtują z kolei „Oczekiwania obywateli wobec administracji”, którym poświęcony jest następny element analizy. Społeczeństwo ma jasno sprecyzowaną wizję: państwo powinno proaktywnie wdrażać algorytmy, aby radykalnie odciążyć mieszkańców z uciążliwej biurokracji, zautomatyzować wypełnianie skomplikowanych formularzy i drastycznie przyspieszyć procesy decyzyjne. Innowacje te napotykają jednak na wyraźne, wyznaczone przez obywateli „czerwone linie”. Bez-

względny wymogi i warunkiem zaufania pozostaje zasada nadrzędności człowieka (human oversight) – obywatele kategorycznie sprzeciwiają się, aby maszyny samodzielnie podejmowały wiążące decyzje, chociażby w zakresie przyznawania świadczeń społecznych czy rekrutacji edukacyjnych.

Klamrą spinającą ten przekrojowy obraz będą ostateczne „Wnioski: obywatel jako punkt odniesienia”. Udowodniono w nich, że implementacja sztucznej inteligencji w aparacie państwowym nie jest już tylko opcją czy wizerunkowym eksperymentem, lecz palącą koniecznością ekonomiczną

i demograficzną. Aby jednak obywatele obdarzyli cyfrowe państwo zaufaniem, ta ewolucja musi opierać się na w pełni bezpiecznych, suwerennych i lokalnych modelach sztucznej inteligencji – takich jak wielofunkcyjny polski model PLLuM – oraz operować w żelaznych, transparentnych ramach regulacyjnych wyznaczonych przez europejski AI Act i dyrektywę NIS 2. Tylko synergia tych działań zagwarantuje stworzenie nowoczesnego, wydajnego państwa, które w centrum swojej uwagi zawsze stawia bezpieczeństwo i godność człowieka.



Obywatel w świecie nowych technologii

Polska znajduje się obecnie w fascynującym, lecz jednocześnie niezwykle wymagającym momencie historycznym, który przez analityków określany jest mianem „wymuszonej dojrzałości cyfrowej”. W ciągu zaledwie kilku ostatnich lat tradycyjna relacja na linii państwo-obywatel uległa diametralnej, wręcz nieodwracalnej zmianie. Obywatel, niegdyś zmuszony do osobistego, często wielogodzinnego stawiennictwa w urzędach i traktowany nierzadko w kategoriach biernego petenta, dziś z sukcesem ewoluuje w stronę w pełni uprawnionego „klienta cyfrowego”.

Proces powszechnej cyfryzacji usług publicznych w Polsce jest zjawiskiem, które spotkało się z potężnym entuzjazmem społecznym. Jak dowodzą dogłębne badania Polskiego Instytutu Ekonomicznego (PIE), aż 92,5% respondentów kategorycznie deklaruje, że cyfrowe usługi publiczne w sposób realny ułatwiają im codzienne załatwianie spraw urzędowych.⁰¹ Co więcej, aż sześciu na dziesięciu Polaków (60%) uważa z pełnym przekonaniem, że znacznie łatwiej, szybciej i wygodniej jest załatwić sprawę administracyjną przez internet, niż podczas tradycyjnej wizyty w stacjonarnej placówce.⁰² Zaledwie 5,4% badanych przyznaje, że nigdy nie skorzystało z żadnego rozwiązania online oferowanego przez państwo. Wyniki te są bezprecedensowe – dowodzą, że polskie społeczeństwo wykazuje niezwykłą elastyczność i chęć adaptacji do nowych, cyfrowych kanałów komunikacji. Spektakularny sukces platform takich jak mObywatel, e-Urząd czy Internetowe Konto Pacjenta udowadnia, że innowacje technologiczne, o ile są użyteczne, stabilne i intuicyjne dla końcowego odbiorcy, spotykają się z powszechną aprobatą.

Nie można jednak w tej analizie ignorować faktu, że obecny krajobraz technologiczny w Polsce charakteryzuje się bardzo poważną asymetrią kompetencyjną. Z jednej strony państwo dysponuje europejskiej klasy infrastrukturą informatyczną i dynamicznie rozwija-

janymi platformami, z drugiej zaś – równolegle zmaga się z wyraźnym deficytem ponadpodstawowych umiejętności cyfrowych wśród dużej części dorosłych obywateli. Zgromadzone dane badawcze jednoznacznie dowodzą, że to właśnie wiek pozostaje najsilniejszym czynnikiem różnicującym poziom adaptacji do nowych technologii. Najmłodsze pokolenie Polaków (15–24 lata) dorasta w środowisku nierozdzielnie związanym z cyfryzacją, wykazując najwyższy poziom optymizmu i naturalnej płynności w nawigowaniu po świecie algorytmów.

Z kolei w starszych grupach wiekowych narasta wyraźny lęk przed technologicznym wykluczeniem i obawa przed bezpowrotnym zastąpieniem żywego, empatycznego kontaktu z drugim człowiekiem przez bezosobowe, zautomatyzowane systemy maszynowe. Z tego powodu nowa „Strategia Cyfryzacji Polski do 2035 roku” zakłada, że w perspektywie najbliższej dekady przynajmniej podstawowe kompetencje cyfrowe musi posiadać aż 85% obywateli. Postawy te pokazują jednoznacznie: polskie społeczeństwo jest w pełni gotowe na kolejny technologiczny skok w postaci sztucznej inteligencji, jednak wymaga on mądrego, inkluzywnego podejścia, które nie pozostawi w cyfrowym niebycie osób gorzej odnajdujących się w wirtualnej rzeczywistości.

AI w codziennym życiu obywateli

Wkroczenie sztucznej inteligencji (AI) do codziennego życia, pracy i edukacji Polaków nastąpiło z dynamiką, która zaskoczyła samych analityków rynku. Polska w bardzo krótkim czasie wyrosła na paradoksalnego lidera w zakresie adaptacji rozwiązań opartych na modelach generatywnych. Według prestiżowego, globalnego raportu opracowanego przez firmę badawczą KPMG („Sztuczna inteligencja w Polsce. Krajobraz pełen paradoksów”), aż 84% badanych w naszym kraju zadeklarowało, że miało kontakt i korzysta ze sztucznej inteligencji w jakiejkolwiek formie – co odczuwalnie przewyższa średnią światową wynoszącą w tym badaniu 80%. Jeszcze bardziej uderzający jest fakt, że regularne i świadome korzystanie z AI deklaruje dziś blisko siedmiu na dziesięciu Polaków (69%), co plasuje polskie społeczeństwo wyżej pod kątem zaawansowania cyfrowego w tym obszarze niż obywateli wielu wysokorozwiniętych państw zachodnich.

Należy jednak postawić kluczowe pytanie: w jaki sposób i za pomocą jakich narzędzi Polacy eksplorują możliwości sztucznej inteligencji? Aż 83% polskich użytkowników AI wykorzystuje ją przede wszystkim w celach prywatnych, a 71% pracujących respondentów stosuje te rozwiązania w swoim środowisku zawodowym. Niestety, w tym entuzjazmie kryje się potężne ryzyko dla bezpieczeństwa informacji. Zdecydowana większość, bo aż 90% pytanych użytkowników, w celu wspierania się AI wybiera narzędzia ogólnodostępne i publiczne, a w przypadku 77% respondentów wybór ten pada wyłącznie na rozwiązania całkowicie bezpłatne. Najpopularniejszymi narzędziami pozostają darmowe generatory tekstu, tłumacze oparte na sieciach neuronowych oraz boty konwersacyjne (np. darmowe wersje ChatGPT, Copilot). Masowe i niekontrolowane powierzanie tym otwartym modelom informacji prywatnych, a nierzadko również poufnych danych firmowych czy urzędowych (tzw. zjawisko „shadow AI”), odbywa się najczęściej bez jakiejkolwiek refleksji nad polityką prywatności i faktem, że dane te posłużą globalnym korporacjom technologicznym (Big Tech) do dalszego trenowania ich algorytmów.

Polska adopcja sztucznej inteligencji to – jak trafnie ujmują to badacze – „krajobraz pełen paradoksów”. Pomimo niezwykle wysokich wskaźników codziennego użytkowania

technologii, świadomość uregulowań prawnych oraz realnych zagrożeń w polskim społeczeństwie pozostaje na alarmująco niskim poziomie. Zaledwie 29% Polaków ukończyło jakiejkolwiek szkolenie (nawet nieformalne) z zakresu sztucznej inteligencji, a aż 90% respondentów nie ma wręcz pojęcia o istnieniu jakichkolwiek uregulowań prawnych dotyczących tej technologii (np. europejskiego AI Act).

Brak fundamentalnej edukacji przekłada się na gigantyczny deficyt zaufania. Zaledwie 41% Polaków deklaruje, że ufa sztucznej inteligencji, co jest wynikiem zauważalnie niższym niż średnia globalna (46%). Z badań wynika, że aż 55% Polaków odczuwa wyraźne obawy względem AI, dostrzegając w niej więcej zagrożeń niż korzyści. Czego dokładnie obawia się społeczeństwo zanurzone w cyfrowym świecie? Na pierwszy plan zdecydowanie wysuwają się cyberzagrożenia – aż 89% Polaków wskazuje na wysokie ryzyko ataków hakerskich z użyciem AI. Ponadto 54% badanych obawia się masowej dezinformacji i zjawiska deepfake (strach przed manipulacjami wyborczymi i politycznymi), a znacząca część społeczeństwa paraliżuje lęk przed utratą prywatności, kradzieżą tożsamości i całkowitą utratą kontroli nad własnymi danymi osobowymi. Obywatele intuicyjnie czują potęgę algorytmów, intensywnie z nich korzystają, ale brakuje im systemowej tarczy ochronnej.

Oczekiwania obywateli wobec administracji

Znając złożony, cyfrowy profil polskiego obywatela, jego powszechne nawyki w korzystaniu z darmowych narzędzi GenAI oraz głęboki stosunek pełen obaw, należy postawić kluczowe dla państwa pytanie: jak społeczeństwo widzi rolę sztucznej inteligencji w architekturze administracji publicznej? Wyniki szeroko zakrojonych sondaży dostarczają jednoznacznej odpowiedzi: obywatele nie tylko są gotowi na wdrożenie AI do polskich urzędów, ale wręcz mają w stosunku do niej wysoce sprecyzowane, bardzo pragmatyczne oczekiwania.

Aż 67% Polaków opowiada się za szerszym stosowaniem sztucznej inteligencji w administracji publicznej. Co więcej, sześciu na dziesięciu ankietowanych (60,4%) jest mocno przekonanych, że to właśnie państwo powinno aktywnie i celowo wykorzystywać nowoczesne algorytmy przy projektowaniu i tworzeniu cyfrowych usług dla obywateli. Obywatele są wyraźnie zmęczeni uciążliwą, przedłużającą się biurokracją i skomplikowanym, hermetycznym językiem prawnym urzędów. Jak wynika z analiz PIE, niemal co trzeci obywatel (32%) pokłada w sztucznej inteligencji ogromne nadzieje na przyspieszenie procedur, bezbłędne zautomatyzowanie procesów i drastyczne skrócenie czasu oczekiwania na wydanie decyzji urzędowych. Społeczeństwo wprost oczekuje wdrożenia asystentów głosowych, inteligentnych wyszukiwarek dokumentów oraz zaawansowanych chatbotów (np. integracji wielofunkcyjnego modelu PLLuM w aplikacji mObywatel), które uproszczą interakcję z państwem.

Aby precyzyjnie zrozumieć mierzalną wartość, jaką obywatele przypisują innowacjom w urzędach, badacze PIE przeprowadzili unikalny, zaawansowany eksperyment ankietowy oparty na metodologii Discrete Choice Experiment (DCE). Wyniki okazały się przełomowe – obywatele potrafili bardzo dokładnie „wycenić” własny komfort i czas. Za ewentualną, dodatkową funkcjonalność publicznej aplikacji, która za pomocą sztucznej inteligencji odpowiadałaby na zawile pytania

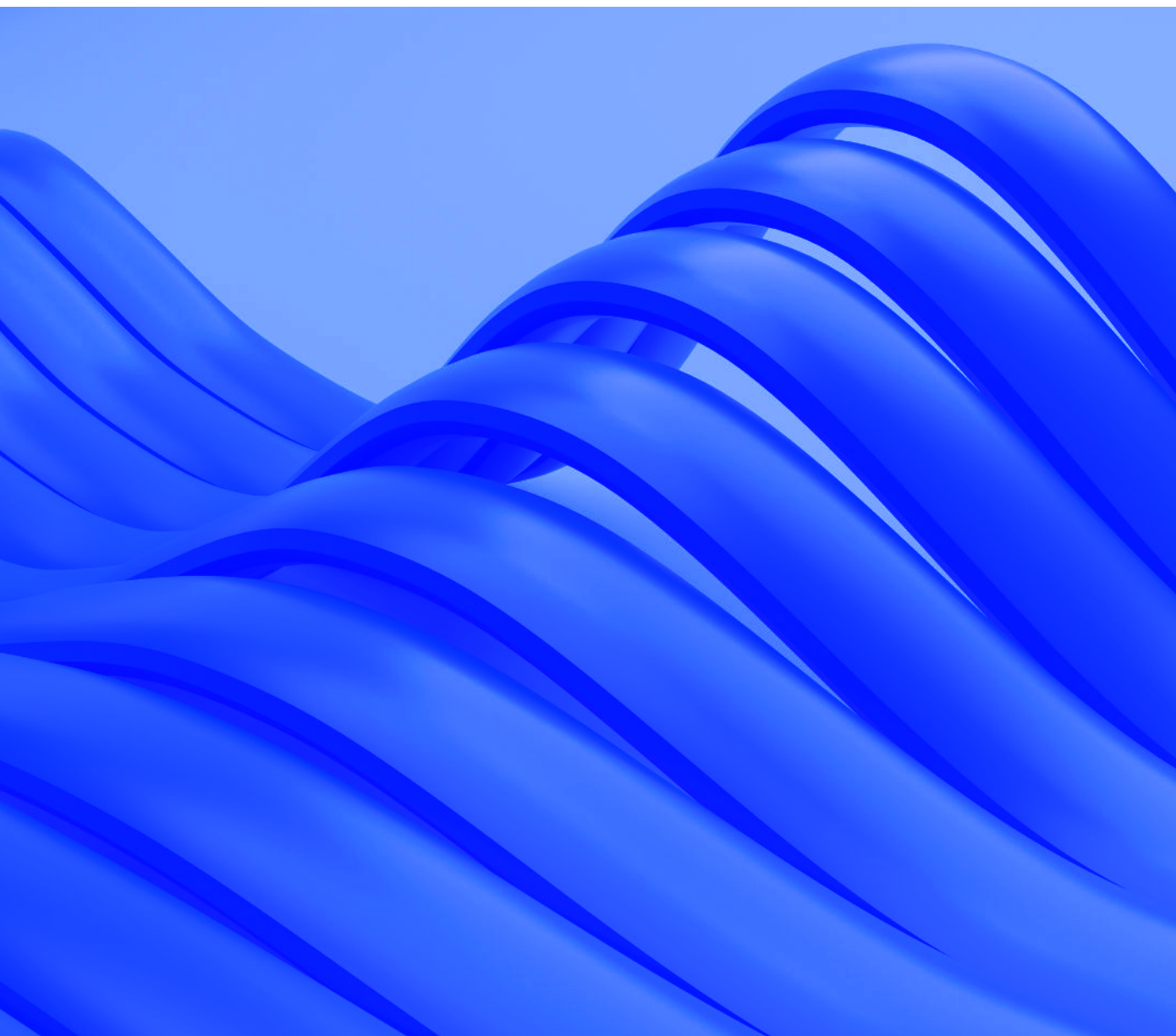
urzędowe w języku potocznym (naturalnym), obywatele byliby skłonni zapłacić 2,53 PLN miesięcznie. Zdecydowanie najwyższą wartość przypisano jednak rozwiązaniom całkowicie wyręczającym człowieka z biurokracji – aplikację, która z wykorzystaniem AI potrafiłaby samodzielnie i bezbłędnie wstępnie wypełniać skomplikowane formularze do danego urzędu, obywatele wycenili aż na 6,02 PLN miesięcznie. Eksperyment ten obnaża potężną, społeczną potrzebę uwolnienia się od urzędowych formalności za pomocą inteligentnych asystentów.

Zgoda na wdrażanie algorytmów przez aparat państwowy nie jest jednak w żadnym wypadku bezwarunkowa. Polacy wyznaczają tu bardzo ostre, jaskrawe „czerwone linie”, których administracja publiczna nie ma prawa przekraczać. Zdecydowana większość społeczeństwa kategorycznie sprzeciwia się temu, aby sztuczna inteligencja w pełni samodzielnie podejmowała kluczowe decyzje mające istotny wpływ na ludzkie życie. Państwo nie posiada społecznego mandatu na wykorzystywanie algorytmów m.in. do ostatecznego orzekania o przyznaniu świadczeń socjalnych (np. 800+, zasiłków, rent), czy też do zautomatyzowanego podejmowania decyzji o rekrutacji do szkół i na uczelnie. Absolutnym fundamentem akceptacji innowacji jest zasada nadrzędności i nadzoru człowieka (human oversight). Z poglądem tym zgadza się zresztą 91% samych urzędników, uważających, że to człowiek, a nie maszyna, musi

ponosić ostateczną, prawną i moralną odpowiedzialność za wydaną decyzję, a system AI może być jedynie narzędziem wspierającym analitykę.

Potężnym wyzwaniem rzutującym na oczekiwania obywateli pozostaje kwestia polityki prywatności i zaufania instytucjonalnego. Aż 88% ankietowanych stanowczo domaga się, aby obywatele mieli bezwzględnie większą, przejrzystą kontrolę nad tym, jak ich newralgiczne dane (zdrowotne, podatkowe, majątkowe) są przetwarzane przez państwo i algorytmy. Co niezwykle niepokojące, ad-

ministracja szczebla centralnego cieszy się obecnie u obywateli bardzo niskim poziomem zaufania w kontekście odpowiedzialnego tworzenia AI. Obywatele znacznie chętniej (ok. 30% wskazań) ufają uniwersytetom publicznym i instytutom badawczym, uznając, że to one tworzą technologie w interesie społecznym. Zaufanie do rządu i administracji w tym zakresie deklaruje zaledwie od 11% do 15% respondentów. Odbudowa tego zaufania – poprzez pełną transparentność, otwarte modele i rygorystyczne bezpieczeństwo – stanowi fundamentalne zadanie państwa w procesie cyfryzacji na lata 2026-2029.



Wnioski: Obywatel jako punkt odniesienia

Głęboka analiza aktualnego stanu cyfryzacji społeczeństwa, nawyków technologicznych Polak i Polaków oraz ich twardych, sprecyzowanych oczekiwań względem państwa, prowadzi do jednoznacznych, strukturalnych konkluzji, które determinują przyszłość polskiej administracji:

Powszechna gotowość do cyfrowej ewolucji i oczekiwanie zmian: Społeczeństwo obywatelskie w Polsce jest pragmatycznie i w pełni gotowe na gruntowną implementację rozwiązań opartych na sztucznej inteligencji. Sukcesy w adaptacji platform cyfrowych takich jak mObywatel udowadniają ponad wszelką wątpliwość, że Polacy pragną nowoczesnego państwa. Oczekują oni odciążenia od biurokracji, błyskawicznego wsparcia (asystenci AI odpowiadający w języku potocznym) oraz automatycznego wypełniania skomplikowanych formularzy urzędowych.

Demograficzna i ekonomiczna konieczność działania: Decyzja o wdrażaniu systemów AI do administracji publicznej nie jest już wyłącznie wyborem wizerunkowym czy eksperymentem technologicznym – to absolutna i paląca konieczność strukturalna i ekonomiczna. Obiektywne wskaźniki dowodzą, że starzejące się w ogromnym tempie społeczeństwo oraz drastycznie kurczący się rynek pracy wymuszają radykalną optymalizację sektora publicznego. Zaawansowane analizy ekonomiczne wskazują, że powszechne zastosowanie generatywnej AI w administracji oraz ograniczenie dzięki temu biurokratycznych ciężarów może przynieść polskiej gospodarce gigantyczne oszczędności rzędu od 6 do 7 miliardów złotych rocznie.

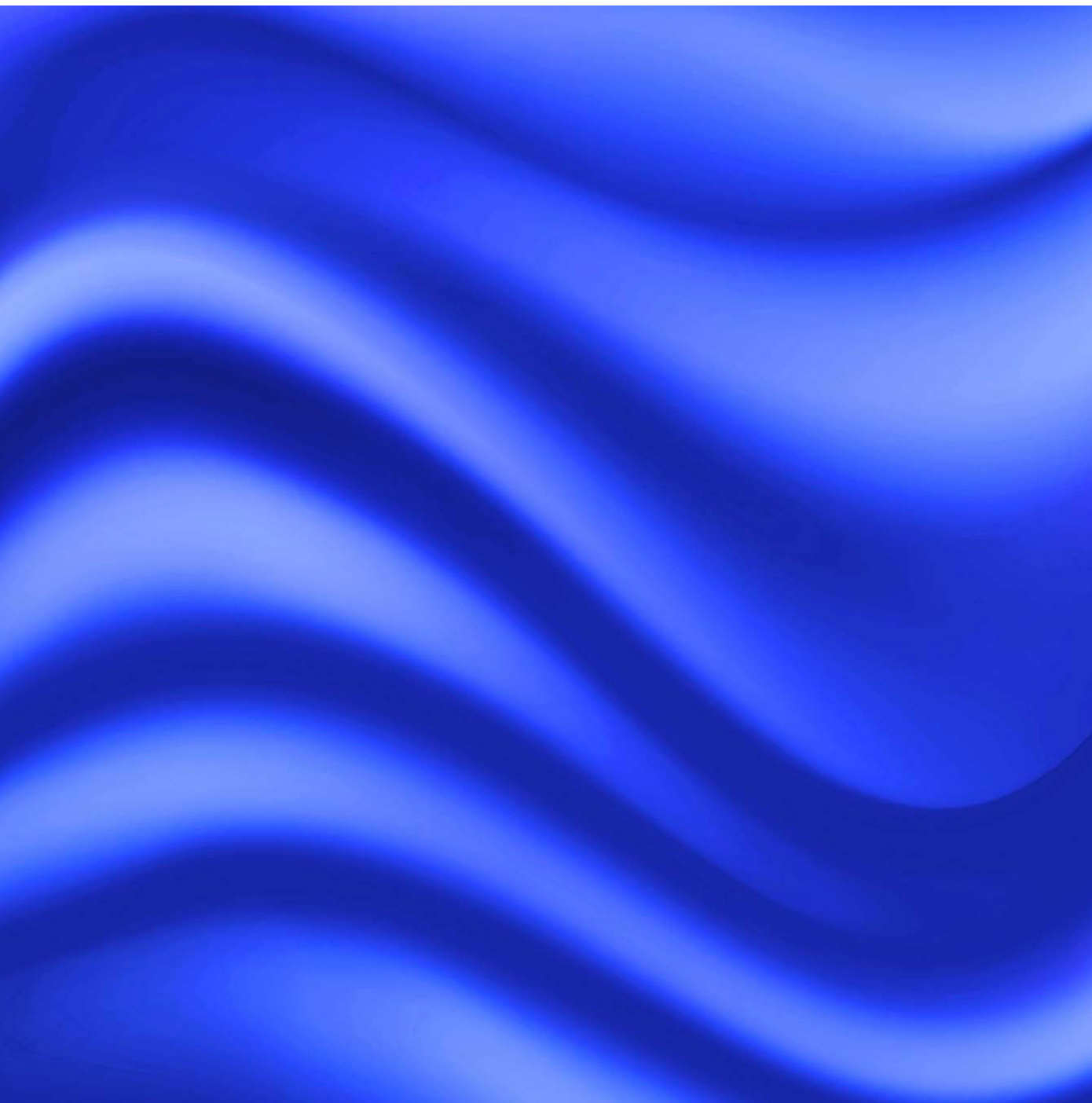
Potrzeba bezpiecznych, suwerennych i lokalnych modeli AI: Biorąc pod uwagę alarmujące dane o tym, że aż 89% Polaków odczuwa głęboki strach przed cyberatakami, a prywatność danych jest ich główną obawą, polskie państwo nie może oprzeć swojej krytycznej

architektury i relacji z obywatelem na publicznych, zagranicznych usługach chmurowych (zjawisko shadow AI). Obywatele nie ufają globalnym korporacjom w kwestii wrażliwych danych. Fundamentem budowy bezpiecznego państwa musi być cyfrowa suwerenność. Zapewnią ją wyłącznie państwowe chmury obliczeniowe (suwerenna chmura) ze ścisłą kontrolą nad danymi oraz bezpieczne, otwarte (open-source), lokalne modele sztucznej inteligencji. Doskonałym przykładem jest wspierany przez instytuty badawcze narodowy model wielofunkcyjny PLLuM czy społecznościowy model Bielik. Rozwiązania te gwarantują utrzymanie pełnej tajemnicy państwowej i kulturowego, polskiego kontekstu.

Żelazne ramy bezpieczeństwa w zgodzie z AI Act i dyrektywą NIS 2: Polacy wprost (63,2% ankietowanych) domagają się silnych państwowych regulacji chroniących ich przed nadużyciami algorytmów. Administracja musi stanowić tu wzór do naśladowania. Oznacza to bezwzględną, restrykcyjną implementację zapisów Aktu w sprawie sztucznej inteligencji (AI Act) – w tym przede wszystkim zagwarantowanie nadzoru wykwalifikowanego człowieka (human-in-the-loop) nad każdą decyzją urzędową wygenerowaną przez maszynę oraz pełną transparentność procesów. Równocześnie niezbędne jest stosowanie rygorów dyrektywy NIS 2 (oraz znowelizowanej polskiej ustawy o KSC), które wymuszają na kadrze kierowniczej urzędów tworzenie procedur reagowania na cyberataki i ochronę przed wyciekiem wrażliwych baz danych obywateli.

Ostateczny wniosek raportu jest czytelny: obywatel jest i musi pozostać nadrzędnym punktem odniesienia dla projektowanego systemu państwowego. Państwo nie może wdrażać najnowszych technologii wyłącznie w celu sztucznego zawyżania statystyk innowacyjności czy cyfryzacji na arenie europejskiej. Każdy nowo powołany do życia publiczny algorytm musi zostać krytycznie

oceniony pod kątem jego realnego, wymiernego wpływu na jakość życia jednostki, przy absolutnym poszanowaniu jej prywatności, godności i nienaruszalnego prawa do obsługi przez człowieka w sytuacjach kryzysowych. Tylko realizacja takich założeń zapewni bezpieczne i suwerenne państwo, realnie odpowiadające na potrzeby społeczeństwa ery sztucznej inteligencji.



Rozdział 1

Obywatel w świecie nowych technologii

1. Polska pozycja (PIE): Polski Instytut Ekonomiczny, 67% Polaków jest za szerszym stosowaniem AI w administracji publicznej, <https://pie.net.pl/67-proc-polakow-jest-za-szerszym-stosowaniem-ai-w-administracji-publicznej/>, dostęp: 25.03.2026.

2. Ibidem.

02

Kontekst technologiczny



krajowa
izba
gospodarki
cyfrowej

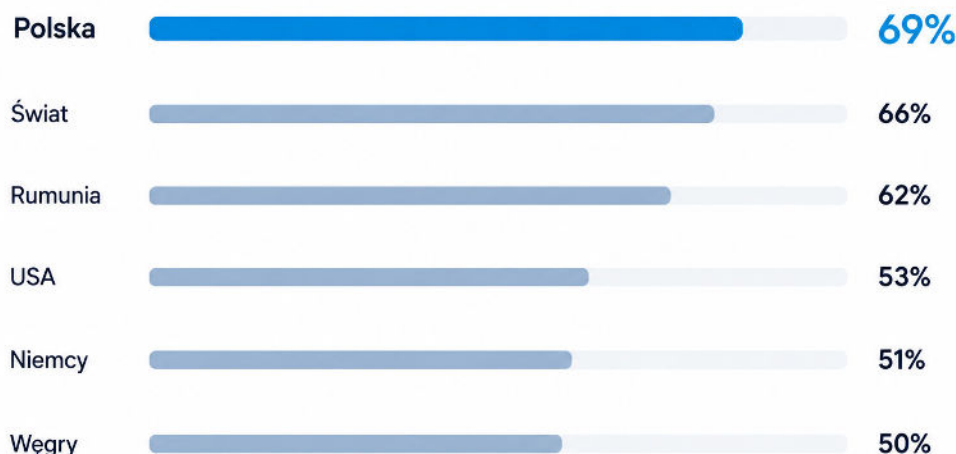
AI i transformacja pracy

Obraz rynku pracy i dynamika adopcji sztucznej inteligencji

Polska powoli wyrasta na jednego z liderów globalnej adaptacji sztucznej inteligencji, wykazując poziom wykorzystania technologii wyższy niż średnia światowa oraz wyniki odnotowane w krajach takich jak USA czy Niemcy. Podczas gdy globalnie do regularnego korzystania z AI przyznaje się 66% respondentów, w Polsce wskaźnik ten wynosi 69%. Skala kontaktu z tą technologią jest jeszcze szersza – 84% Polaków deklaruje, że korzystało z rozwiązań opartych na AI w jakiegokolwiek formie.

Wykorzystanie AI w sposób regularny

(% użytkowników)



Wykres 2.1. Źródło: Raport KPMG w Polsce „Sztuczna inteligencja w Polsce. Krajobraz pełen paradoksów” strona nr 8

Wdrażanie technologii opartych na sztucznej inteligencji na polskim rynku pracy wykazuje wyraźną dynamikę wzrostu. Zestawienie danych z lat 2024-2025 wskazuje na utrwalanie się roli tych rozwiązań w sferze zawodowej. Według raportu opracowanego przez NASK i Międzynarodową Organizację Pracy (ILO) pod koniec 2024 roku, 30,5% pracujących Polaków deklarowało korzystanie z narzędzi generatywnej sztucznej inteligencji (GenAI)

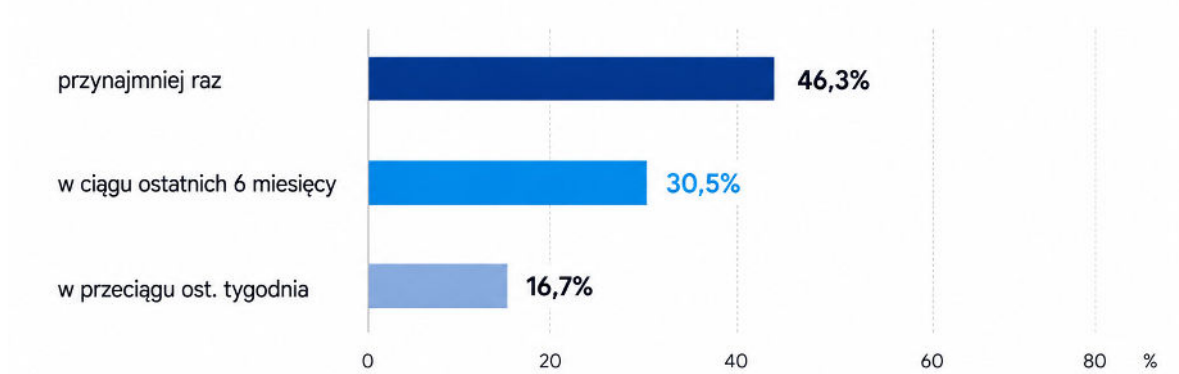
w pracy w ciągu ostatnich sześciu miesięcy, a przynajmniej raz miało z nimi kontakt 46,3% osób.

Nowsze dane pochodzące z raportu KPMG „Sztuczna inteligencja w Polsce. Krajobraz pełen paradoksów” wskazują na dalsze upowszechnianie się tej technologii. Obecnie 54% osób pracujących w Polsce deklaruje,

że korzysta z AI w celach zawodowych intencjonalnie i regularnie (co najmniej raz na kilka miesięcy). Warto zauważyć, że proces ten często odbywa się oddolnie: prawie 90% użytkowników wykorzystuje w pracy publicz-

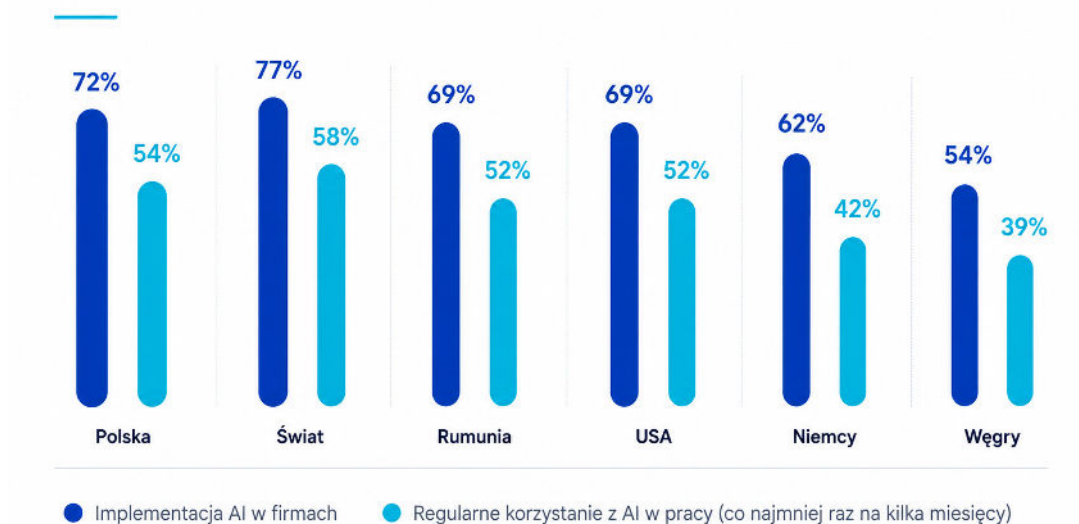
nie dostępne narzędzia, a 77% sięga po ich bezpłatne wersje. Jednocześnie oficjalne wdrożenie AI deklaruje co najmniej 72% organizacji w Polsce.⁰¹

Korzystanie z Generatywnej AI w różnych przedziałach czasowych



Wykres 2.2 Źródło: Raport badawczy „Generatywna sztuczna inteligencja a polski rynek pracy” strona nr 36

Wykorzystanie i wdrażanie AI w pracy w wybranych krajach



Wykres 2.3 Źródło: Raport KPMG w Polsce „Sztuczna inteligencja w Polsce. Krajobraz pełen paradoksów”

Zjawisko Shadow AI i wyzwania kompetencyjne

Upowszechnienie sztucznej inteligencji wyprzedza procesy regulacyjne wewnątrz organizacji, co prowadzi do powstawania zjawiska tzw. „Shadow AI” – korzystania z technologii poza oficjalnymi procedurami pracodawcy. Dane z 2025 roku wskazywały, że 68,1% zatrudnionych Polaków nie otrzymało od swoich pracodawców żadnych wytycznych dotyczących zasad korzystania z GenAI. Problem ten dotyczył również administracji publicznej, gdzie według raportu NASK „AI w administracji publicznej” aż 70% badanych urzędników wskazywało na brak wewnętrznych procedur jako główną barierę we wdrażaniu rozwiązań AI.

Dane z raportu KPMG „Sztuczna inteligencja w Polsce - Krajobraz pełen paradoksów” wskazują, że deficyt uregulowań utrzymuje się: 7 na 10 Polaków korzysta z AI bez przeszkolenia i często bez wiedzy pracodawcy. Ponad połowa (55%) pracowników używających sztucznej inteligencji nie ujawnia tego faktu, prezentując wygenerowane wyniki jako własną pracę. Równocześnie aż 90% badanych z Polski nie ma świadomości istnienia obowiązujących regulacji dotyczących AI, a co trzeci pracownik wykorzystujący AI w pracy przyznaje się do działań sprzecznych z polityką swojej firmy.⁰²

Sytuacja ta jest częściowo wynikiem presji technologicznej. Prawie połowa (48%) pracowników obawia się, że jeśli nie będą uży-

wać AI, pozostaną w tyle za zmianami na rynku pracy. Widać również wyraźną lukę kompetencyjną: choć 60% ankietowanych twierdzi, że potrafi efektywnie posługiwać się sztuczną inteligencją, w rzeczywistości jedynie 29% przeszło jakiegokolwiek szkolenie w tym zakresie.

Oczekiwane korzyści i automatyzacja zadań w administracji publicznej

Pomimo zidentyfikowanych ryzyk, poziom akceptacji dla obecności AI w miejscu pracy w Polsce wynosi 77%. Z danych KPMG wynika, że 76% użytkowników zauważyło poprawę dostępności zasobów i wydajności dzięki automatyzacji powtarzalnych zadań. Ogólny wzrost efektywności pracy potwierdza 64% badanych, chociaż dla 1/3 oznacza to również zwiększone obciążenie obowiązkami wynikające z wdrożenia nowych technologii.

W sektorze administracji publicznej pracownicy dostrzegają w sztucznej inteligencji przede wszystkim narzędzie do optymalizacji procesów. Aż 81% ankietowanych urzędników wskazuje, że AI wesprze automatyzację rutynowych zadań, a 76% uważa, że pozwoli to na wymierne oszczędności czasu. Badania NASK pokazują, że urzędnicy widzą największy potencjał wsparcia w konkretnych działaniach: tłumaczeniach oraz redakcji treści w językach obcych (77%), pracy z dokumentami w różnych formatach (72%) oraz błyskawicznym wyszukiwaniu informacji w dokumentach urzędowych przy pomocy pytań zadawanych w języku potocznym (67%).⁰³

W jakim stopniu wymienione funkcjonalności AI mogą wspierać pracę urzędników?

Praca z dokumentami w różnych formatach

(np. możliwość przeszukiwania i analizowania plików PDF, Word, skanów)



Wyszukiwanie informacji w dokumentach urzędowych za pomocą pytań zadawanych w języku potocznym



Wykres 2.4 Źródło: AI w e-administracji publicznej – perspektywa urzędników i instytucji

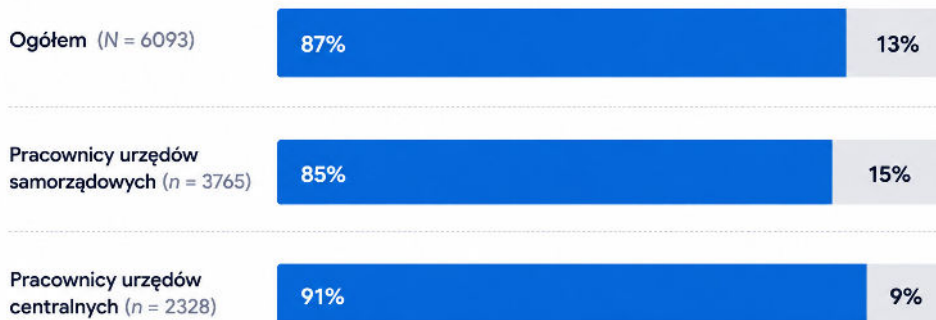
Personel e-administracji wykazuje dużą gotowość do podnoszenia swoich kwalifikacji – aż 87% badanych urzędników deklaruje chęć udziału w dedykowanych szkoleniach z zakresu wykorzystania AI.

Dążenie to pokrywa się z rynkową potrzebą inwestowania w cyfrową edukację kadr, co w przyszłości może przełożyć się na sprawniejszą komunikację i szybszą realizację usług publicznych dla obywateli.

Chęć udziału w szkoleniach z zakresu wykorzystania narzędzi AI w pracy – typ urzędu

■ Zdecydowanie + Raczej tak

■ Zdecydowanie + Raczej nie



Wykres 2.5 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik-listopad 2025, NASK-PIB.

Wszystkie te twarde dane potwierdzają jedną, istotną kwestię. Transformacja rynku pracy już stała się faktem, na ten moment w większości przypadków działa to oddolnie. Teraz prawdziwym zadaniem i odpowiedzialnością

dla całego środowiska odpowiedzialnego za realne wdrożenia i szkolenia, jest masowe przeszkolenie kadr, by wszystko odbyło się zgodnie ze wszystkimi normami i ograniczeniami.⁰⁴

Rozwój AI i modeli językowych

Transformacja technologiczna, będąca stałym elementem nowoczesnych systemów społeczno-gospodarczych, gwałtownie przyspieszyła wraz z rozwojem generatywnej sztucznej inteligencji. Od momentu premiery przełomowych modeli, takich jak ChatGPT pod koniec 2022 roku, obserwujemy bezprecedensowy wzrost dostępności narzędzi, które fundamentalnie zmieniają organizację pracy, procesy produkcyjne oraz interakcje społeczne. Potencjał transformacyjny tej technologii jest obecnie porównywany do wynalezienia elektryczności czy internetu. Choć modele te potrafią tworzyć treści, obrazy i kod na poziomie przewyższającym przeciętnego użytkownika, ich dynamiczny rozwój rodzi naturalne pytania o przyszłość dochodów i stabilność miejsc pracy. Zanim jednak modele generatywne stały się narzędziem masowym, sztuczna inteligencja przechodziła przez fazy laboratoryjnych eksperymentów, teoretycznych sporów oraz technicznych ograniczeń. Zrozumienie dzisiejszej rewolucji wymaga spojrzenia na historię AI jako na proces ciągłego przesuwania granic między wizjonerską koncepcją a fizyczną wydajnością infrastruktury obliczeniowej. Poniżej zaprezentowano w obrazowy sposób, dane dotyczące rozwoju AI, wliczając w to historię oraz główne paradygmaty.

HISTORIA AI

1950

TEST TURINGA

Alan Turing publikuje „Computing Machinery and Intelligence” i wprowadza Test Turinga – pierwszy miernik inteligencji maszynowej.

1956

POCZĄTEK AI

Podczas konferencji w Dartmouth powstaje termin „sztuczna inteligencja” i formalnie rozpoczyna się rozwój dziedziny.

1958

PERCEPTRON

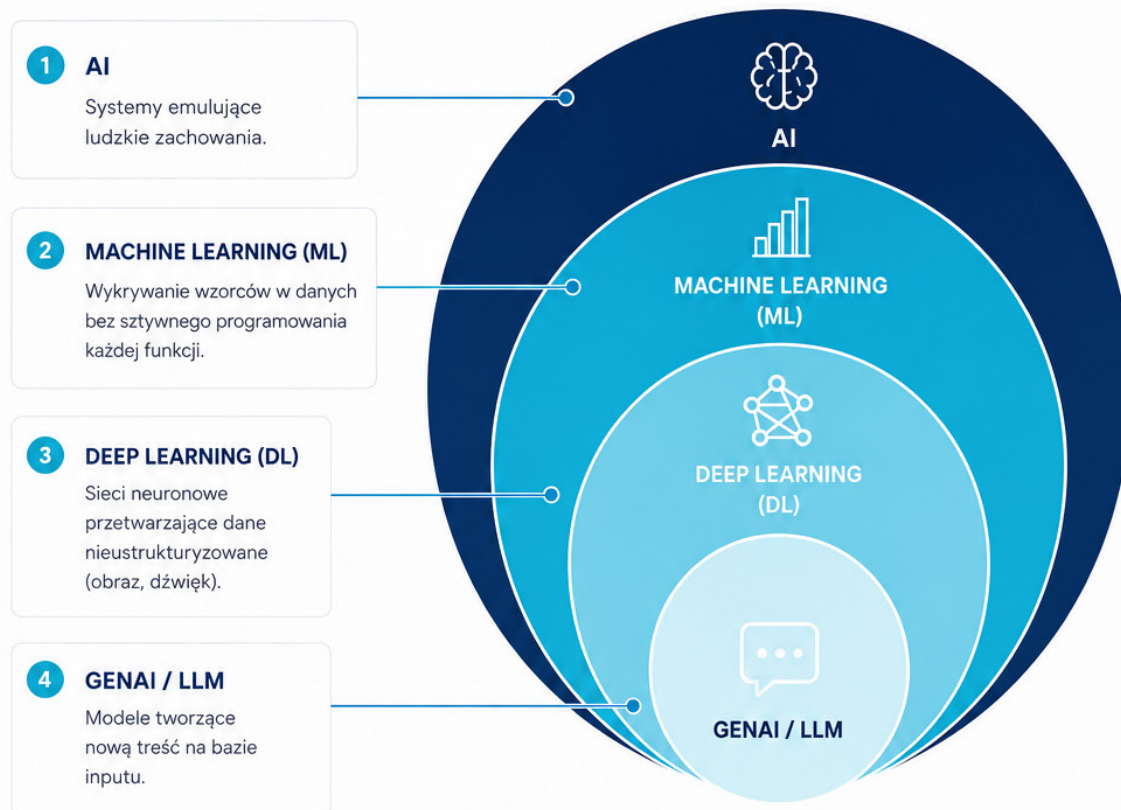
Frank Rosenblatt tworzy pierwszy model sztucznej sieci neuronowej zdolny do uczenia się na podstawie danych wejściowych.



Wykres 2.6 Źródło: Opracowanie własne

Architektura i Doktryny Uczenia

Systemy inteligentne dzielą się na warstwy o rosnącej specjalizacji:



Paradygmaty nauczania

- 1 UCZENIE NADZOROWANE (SUPERVISED)**
Nauka na parach „wejście-wynik” (etykiety).
Stosowane w klasyfikacji i regresji liniowej.
- 2 UCZENIE NIENADZOROWANE (UNSUPERVISED)**
Szukanie ukrytych struktur w surowych danych
(klastrowanie K-means, redukcja PCA).
- 3 UCZENIE PRZEZ WZMACNIANIE (REINFORCEMENT)**
Nauka przez interakcję ze środowiskiem
i system kar/nagród (Q-Learning, DQN).
- 4 RLHF (HUMAN FEEDBACK)**
Klucz do sukcesu ChatGPT. Ludzie oceniają odpowiedzi AI,
tworząc „model nagrody”, który uczy maszynę bezpiecznej
i naturalnej komunikacji.

Wykres 2.7 Źródło: Opracowanie własne

Modele językowe i sieci neuronowe



DUŻE MODELE JĘZYKOWE (LLM)

Oparte na architekturze Transformer, wykorzystują mechanizm Self-Attention, by rozumieć kontekst słów niezależnie od ich odległości w zdaniu.

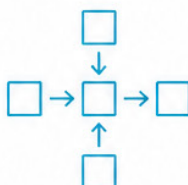
Podział architektoniczny

1 AUTOREGRESYWNE (DECODER-ONLY)



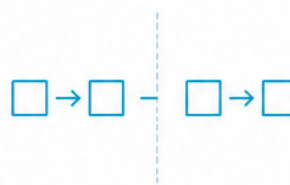
Przewidują kolejne słowo.
Fundament serii GPT i Llama.

2 AUTOENKODERY (ENCODER-ONLY)



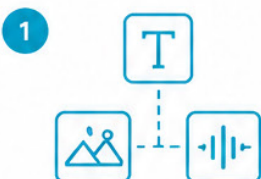
Głębokie rozumienie tekstu
(np. BERT). Idealne do
klasyfikacji i wyszukiwania.

3 HYBRYDOWE (ENCODER-DECODER)



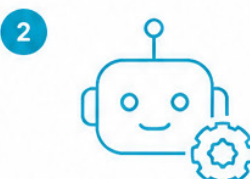
Tłumaczenia i streszczenia
(np. T5).

Nowe trendy



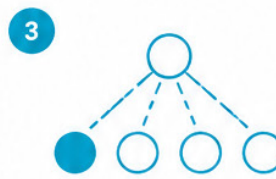
MULTIMODALNOŚĆ

Modele natywnie obsługujące
tekst, obraz i dźwięk
(GPT-4o, Gemini).



AGENTOWOŚĆ

Dedykowane autonomiczne agenty,
które wykonują zlecone zadania
niezależnie, bez akceptacji
użytkownika.



MIXTURE OF EXPERTS (MOE)

Aktywacja tylko fragmentów sieci
„ekspertów” dla konkretnego
zadania, co drastycznie tnie koszty
energii przy zachowaniu skali wiedzy.

Wykres 2.8 Źródło: Opracowanie własne

Typy Uczenia Maszynowego



UCZENIE NADZOROWANE (SUPERVISED LEARNING)

Algorytmy uczą się na podstawie oznakowanych danych wejściowych i znanych wyników.



KLASYFIKACJA

- filtrowanie spamu
- wykrywanie nadużyć
- rozpoznawanie obrazów



REGRESJA

- prognozowanie sprzedaży
- przewidywanie cen
- szacowanie ryzyka



UCZENIE NIENADZOROWANE (UNSUPERVISED LEARNING)

Algorytmy odkrywają ukryte struktury i wzorce w danych bez znanych etykiet.



REDUKCJA WYMIAROWOŚCI

- wizualizacja danych
- kompresja informacji
- eksploracja danych



KLASTROWANIE

- segmentacja klientów
- grupowanie dokumentów
- wykrywanie anomalii



UCZENIE PRZEZ WZMACNIANIE (REINFORCEMENT LEARNING)

Agent uczy się poprzez interakcję ze środowiskiem, maksymalizując nagrodę długoterminową.



- optymalizacja strategii
- zarządzanie zasobami
- robotyka i sterowanie
- rekomendacje sekwencyjne
- gry i symulacje

ZASTOSOWANIE

Uczenie maszynowe znajduje zastosowanie w wielu branżach i obszarach życia, wspierając automatyzację decyzji, analizę danych i optymalizację procesów.



Finanse



Zdrowie



Administracja
publiczna



Przemysł



Handel
i e-commerce



Edukacja

Wykres 2.9 Źródło: Opracowanie własne

Analiza zgromadzonych danych prowadzi do sformułowania kluczowych wniosków dotyczących obecnego etapu transformacji cyfrowej w Polsce:

- **Prym oddolnej transformacji:** Rewolucja technologiczna w polskich przedsiębiorstwach i urzędach ma charakter oddolny i de facto już się dokonała. Fakt, że 54% osób pracujących regularnie korzysta z AI, a niemal 90% z nich robi to przy użyciu publicznych, często bezpłatnych narzędzi, świadczy o tym, że pracownicy samodzielnie wdrażają innowacje, by sprostać presji rynkowej, nie czekając na oficjalne szkolenia czy wytyczne.
- **Krytyczna potrzeba systemowych uregulowań:** Skala zjawiska „Shadow AI” generuje realne ryzyka technologiczne i prawne. Skoro 68,1% zatrudnionych nie otrzymało żadnych zasad korzystania z AI, a ponad połowa użytkowników ukrywa fakt wspierania się sztuczną inteligencją przed przełożonymi, niezbędne jest pilne przejęcie inicjatywy przez kadrę zarządzającą. Brak oficjalnych procedur przy jednoczesnym masowym wykorzystaniu technologii to prosta droga do wycieków

danych i naruszeń bezpieczeństwa cyfrowego.

- **Luka kompetencyjna jako bariera:** Istnieje wyraźny rozdźwięk między deklarowaną biegłością a realnym przygotowaniem kadr. Choć większość użytkowników uważa, że potrafi efektywnie posługiwać się AI, to tylko 29% przeszło jakiejkolwiek merytoryczne przeszkolenie. Bez systemowego wsparcia edukacyjnego, pracownicy będą powielać błędy, halucynacje modeli i niebezpieczne nawyki, co zamiast zwiększać produktywność, może prowadzić do spadku jakości pracy i frustracji.
- **Odpowiedzialność instytucjonalna:** Transformacja rynku pracy stała się faktem i nie można jej już zatrzymać za pomocą zakazów. Teraz prawdziwym zadaniem dla liderów biznesu i administracji publicznej jest wzięcie odpowiedzialności za ten proces poprzez wdrożenie masowych szkoleń i jasnych ram prawnych. Wykorzystując dużą gotowość pracowników do nauki (87% w przypadku urzędników), należy ucywilizować cyfrową rewolucję tak, aby odbywała się ona w sposób bezpieczny, transparentny i zgodny z obowiązującymi normami.

Rozdział 2

Sekcja: AI i transformacja pracy

1. KPMG w Polsce, KPMG AI Trust 2025, pl-Raport-KPMG-w-Polsce-KPMG-AI-Trust-2025-web.pdf, s. 8.
2. Generatywna sztuczna inteligencja a polski rynek pracy, Generatywna-sztuczna-inteligencja-a-polski-rynek-pracy.pdf, s. 36.
3. KPMG w Polsce, KPMG AI Trust 2025, pl-Raport-KPMG-w-Polsce-KPMG-AI-Trust-2025-web.pdf
4. Badanie ilościowe „AI w e-administracji publicznej”, październik-listopad 2025, NASK-PIB.

03

Regulacje i zmiany w prawie



krajowa
izba
gospodarki
cyfrowej

Regulacje i odpowiedzialność

Wprowadzenie

Wdrażanie sztucznej inteligencji do struktur państwowych nie może odbywać się w próżni prawnej, a zaufanie obywateli do nowych technologii buduje się przede wszystkim poprzez przejrzyste zasady i przypisanie odpowiedzialności. Jak pokazują globalne i krajowe badania, krajobraz świadomości prawnej w Polsce w tym zakresie pełen jest jednak intrygujących paradoksów. Z jednej strony, niemal siedmiu na dziesięciu Polaków regularnie korzysta z narzędzi sztucznej inteligencji, co stawia nas powyżej średniej światowej, z drugiej zaś – aż 90% z nich nie ma pojęcia o istnieniu jakichkolwiek regulacji prawnych dotyczących tej technologii⁰¹.

Społeczeństwo intuicyjnie oczekuje jednak konkretnych działań i zabezpieczeń: 70% badanych na świecie domaga się skutecznych przepisów chroniących użytkowników, a w Polsce szczególne obawy budzi brak kontroli nad cyberzagrożeniami oraz masową dezinformacją generowaną przez AI⁰².

Administracja publiczna znajduje się w samym centrum tej transformacji. Wprowadzanie potężnych algorytmów bez odpowiedniego nadzoru rodzi ogromne ryzyko operacyjne i prawne. Niezwykle groźnym, a jednocześnie powszechnym zjawiskiem, jest tzw. „shadow AI”, czyli korzystanie przez pracowników z ogólnodostępnych, darmowych narzędzi poza oficjalnymi procedurami, bez wiedzy i akceptacji ze strony pracodawców⁰³. Biorąc pod uwagę, że prawie 90% użytkowników AI

w miejscu pracy w Polsce sięga po narzędzia publicznie dostępne, w sektorze publicznym stwarza to bezpośrednie zagrożenie wycieku lub utraty poufności danych obywateli. Zjawisko to uwydatnia pilną potrzebę wdrożenia i stanowczego egzekwowania kompleksowych ram regulacyjnych. Architektura prawna bezpiecznego, cyfrowego państwa musi opierać się na trzech potężnych filarach: nowym Akcie w sprawie sztucznej inteligencji (AI Act), zaostrzonej dyrektywie NIS 2 dotyczącej cyberbezpieczeństwa oraz ugruntowanych przepisach ogólnego rozporządzenia o ochronie danych (RODO). Tylko ich ścisłe, synergiczne stosowanie zagwarantuje, że administracja radykalnie podniesie swoją efektywność, nie narażając przy tym fundamentalnych praw mieszkańców.

Akt w sprawie sztucznej inteligencji (AI Act) - Architektura przyszłości, szanse i wyzwania dla cyfrowego państwa oraz sektorów regulowanych

Cyfrowa rewolucja napędzana przez algorytmy uczenia maszynowego przestała być jedynie wizją z laboratoriów badawczych, stając się fundamentem współczesnej gospodarki i administracji. W odpowiedzi na ten bezprecedensowy rozwój, 1 sierpnia 2024 roku w Unii Europejskiej weszło w życie rozporządzenie w sprawie sztucznej inteligencji, powszechnie znane jako AI Act⁰⁴. Pełne stosowanie większości jego przepisów przewidziano na 2 sierpnia 2026 roku. Jest to pierwsza na świecie tak szeroko zakrojona, horyzontalna próba uregulowania systemów opartych na sztucznej inteligencji, której celem jest zapewnienie bezpieczeństwa, poszanowania praw podstawowych oraz stymulowanie innowacyjności na europejskim rynku cyfrowym⁰⁵.

Czym jest AI Act i na czym opiera się jego filozofia?

Filozofia AI Act opiera się na tzw. podejściu opartym na ryzyku (risk-based approach). Oznacza to, że intensywność regulacji i nałożone obowiązki są wprost proporcjonalne do potencjalnych zagrożeń, jakie dany system AI może generować dla jednostki lub społeczeństwa⁰⁶. Prawo to dzieli systemy sztucznej inteligencji na cztery główne kategorie⁰⁷:

1. Systemy o ryzyku nieakceptowalnym (zakazane): Technologie stwarzające zagrożenie, którego nie da się zmitigować. Zalicza się do nich m.in. systemy oceny obywatelskiej (social scoring), nieukierunkowane pobieranie wizerunków twarzy z internetu do tworzenia baz danych, a także systemy rozpoznawania emocji w miejscu pracy czy w instytucjach edukacyjnych.
2. Systemy wysokiego ryzyka: Technologie mające szczególny wpływ na życie, zdro-

wie i prawa podstawowe. Wymagają one spełnienia rygorystycznych norm przed wprowadzeniem na rynek i w trakcie eksploatacji.

3. Systemy o ograniczonym ryzyku: Wymagają przede wszystkim przejrzystości, np. informowania użytkowników, że wchodzi w interakcję z maszyną (chatboty) lub że treści, takie jak obrazy czy wideo (deep-fakes), zostały wygenerowane sztucznie.
4. Systemy o minimalnym ryzyku: Większość dostępnych dziś algorytmów (np. filtry antyspamowe), które mogą być rozwijane i używane swobodnie, z zaleceniem stosowania dobrowolnych kodeksów dobrych praktyk.

Rozporządzenie wprowadza również osobne ramy dla modeli sztucznej inteligencji ogólnego przeznaczenia (GPAI), takich jak potężne modele fundamentowe, nałożając na ich twórców obowiązki w zakresie dokumentacji technicznej, przejrzystości praw autorskich oraz zarządzania ryzykiem systemowym⁰⁸.

Szanse: Zaufanie, innowacje i cyfrowa suwerenność

AI Act to nie tylko zbiór zakazów, ale przede wszystkim narzędzie budowania przewagi konkurencyjnej opartej na zaufaniu. Zharmonizowane przepisy eliminują ryzyko fragmentacji rynku wewnętrznego Unii Europejskiej, zapewniając pewność prawną i równe zasady gry dla przedsiębiorstw operujących na terenie całej wspólnoty.

Dla innowatorów kluczową szansą są tzw. piaskownice regulacyjne (regulatory sandboxes). Państwa członkowskie mają obowiązek ustanowić bezpieczne, kontrolowane środowiska, w których startupy, małe i średnie przedsiębiorstwa (MŚP) oraz jednostki badawcze mogą testować swoje systemy AI pod nadzorem regulatora, jeszcze przed ich wprowadzeniem na rynek⁰⁹. Udział w takich piaskownicach, darmowy dla MŚP, minimalizuje ryzyko kar, ułatwia proces oceny zgodności i przyspiesza transfer technologii z laboratoriów do gospodarki. Dokument wprowadza tzw. „efekt Brukseli”, dążąc do tego, by europejskie standardy etycznej i bezpiecznej sztucznej inteligencji stały się globalnym punktem odniesienia¹⁰.

Zagrożenia i wyzwania: Koszty, biurokracja i odpowiedzialność

Największym wyzwaniem dla podmiotów rynkowych są koszty dostosowania się do nowych przepisów (compliance). Twórcy i podmioty stosujące systemy wysokiego ryzyka muszą wdrożyć rozbudowane systemy zarządzania jakością, dbać o reprezentatywność i czystość danych treningowych w celu unikania dyskryminacji, a także tworzyć zaawansowaną dokumentację techniczną i systemy rejestrowania zdarzeń (logowania)¹¹.

Nieprzestrzeganie przepisów wiąże się z drastycznymi konsekwencjami finansowymi. System kar przewidziany w Akcie operuje

na kwotach jeszcze wyższych niż te znane z RODO. Za stosowanie praktyk zakazanych grożą kary do 35 milionów euro lub do 7% rocznego światowego obrotu przedsiębiorstwa, a za niespełnienie wymogów dla systemów wysokiego ryzyka – do 15 milionów euro lub 3% obrotu¹².

Poważnym ryzykiem na poziomie operacyjnym jest zjawisko „shadow AI”. Niemal 90% osób wykorzystujących AI w Polsce robi to za pomocą ogólnodostępnych, darmowych narzędzi, często bez zgody i wiedzy pracodawcy. Wprowadzanie poufnych danych firmy, tajemnic przedsiębiorstwa czy danych obywateli do publicznych modeli językowych rodzi bardzo wysokie ryzyko wycieku i łamania przepisów, stanowiąc bezpośrednie zagrożenie dla organizacji¹³.

Perspektywa sektora regulowanego i administracji publicznej

Zarówno dla urzędów, jak i dla sektorów silnie uregulowanych (np. bankowość, ubezpieczenia, medycyna), AI Act stanowi prawdziwe trzęsienie ziemi, ponieważ większość wykorzystywanych tam zaawansowanych algorytmów klasyfikowana jest jako systemy „wysokiego ryzyka”.

W usługach finansowych za wysokie ryzyko uznaje się m.in. narzędzia do oceny zdolności kredytowej obywateli czy szacowania ryzyka w ubezpieczeniach na życie i zdrowotnych. W opiece zdrowotnej są to systemy diagnozujące choroby i planujące terapie. W administracji publicznej i wymiarze sprawiedliwości wymogami wysokiego ryzyka objęto algorytmy oceniające kwalifikowalność do świadczeń socjalnych, systemy wspierające organy ścigania, zarządzanie migracją i azylem, a także narzędzia służące do analizy faktów i przepisów na potrzeby sądów¹⁴.

Dla podmiotów publicznych i prywatnych pełniących usługi użyteczności publicznej

wdrożenie takiego systemu wiąże się z koniecznością przeprowadzenia wnikliwej oceny skutków dla praw podstawowych (Fundamental Rights Impact Assessment – FRIA). Analiza ta musi zidentyfikować, jak algorytm wpłynie na poszczególne grupy społeczne i jakie środki zaradcze zostaną podjęte w razie nieprawidłowości. Systemy te muszą być również wprowadzane do ogólnounijnej, publicznej bazy danych.

Kolejnym żelaznym wymogiem z perspektywy administracji jest zapewnienie „nadzoru człowieka” (human oversight) nad maszyną. Architektura algorytmiczna urzędów i banków musi być zaprojektowana tak, aby człowiek mógł w pełni kontrolować maszynę, zrozumieć jej wyniki i w razie konieczności zainterweniować – zignorować podpowiedź systemu, a nawet go zatrzymać. Jest to ściśle powiązane z wymogami przejrzystości: obywatel ma prawo wiedzieć, że rozmawia ze sztuczną inteligencją oraz ma prawo do merytorycznego wyjaśnienia, jaką rolę AI odegrała w podjęciu decyzji w jego sprawie. Wymogi te krzyżują się ściśle z wytycznymi RODO (m.in. zasadą minimalizacji danych i ochroną prywatności w fazie projektowania) oraz dyrektywą NIS 2, która czyni kierownictwo urzędów i spółek bezpośrednio odpowiedzialnym za cyberbezpieczeństwo infrastruktury¹⁵.

Zgodnie z art. 26 ust. 7 AI Act, pracodawcy (w tym jednostki samorządu terytorialnego i administracja) mają bezpośredni obowiązek informowania przedstawicieli pracowników oraz samych pracowników o tym, że będą oni podlegali działaniu systemów wysokiego ryzyka (np. w rekrutacji czy ocenach pracowniczych), jeszcze przed ich wdrożeniem. W praktyce wzmacnia to rolę dialogu społecznego przy transformacji cyfrowej. Jednocześnie, zgodnie z art. 6 ust. 3 AI Act, systemy wymienione w załączniku III nie muszą być uznane za systemy wysokiego ryzyka, jeżeli nie stwarzają znaczącego ryzyka dla praw podstawowych (np. wykonują jedynie wąsko określone zadania przygotowawcze lub pomocnicze). W takim przypadku dostawca systemu ma jednak obowiązek udokumentowania tej oceny i zarejestrowania systemu w publicznej unijnej bazie danych przed wprowadzeniem go do obrotu.

Z tej perspektywy jednym z preferowanych z punktu widzenia zgodności i bezpieczeństwa kierunków dla administracji państwowej staje się suwerenność technologiczna. Inwestycje w lokalne, państwowe chmury obliczeniowe i otwarte, krajowe modele bazowe – takie jak polski projekt wielofunkcyjnego modelu PLLuM, integrowany z aplikacją mObywatel – mogą pozwalać na automatyzację pracy urzędników bez konieczności wysyłania wrażliwych danych obywateli na serwery zagranicznych dostawców. Dotyczy to również innych rozwiązań działających na infrastrukturze lokalnej, bez wysyłania informacji poufnych do chmury.

Wnioski podsumowujące

Europejski Akt o Sztucznej Inteligencji ustanawia nowy paradygmat funkcjonowania cyfrowego państwa. Wymusza on na administracji publicznej i sektorze regulowanym przejście od fascynacji nowinkami technologicznymi do dojrzałego, opartego na rygorystycznych ramach prawnych zarządzania ryzykiem. Choć wiąże się to ze zwiększonymi obciążeniami dokumentacyjnymi i organizacyjnymi, w ogólnym rozrachunku chroni praworządność, eliminuje uprzedzenia maszyn i buduje zaufanie obywateli. W erze, w której dane i algorytmy decydują o bezpieczeństwie socjalnym, zdrowotnym i prawnym społeczeństwa, AI Act staje się niezbędną instrukcją obsługi innowacji służącej człowiekowi.

Z perspektywy sektora publicznego kluczowe znaczenie mają następujące wnioski:

- AI przestaje być obszarem pilotażowym, a staje się w pełni regulowanym elementem funkcjonowania państwa, wymagającym formalnych procedur, nadzoru i zgodności z przepisami.
- Istotną część zastosowań sztucznej inteligencji w administracji publicznej może być klasyfikowana jako systemy wysokiego ryzyka, co w praktyce oznacza konieczność każdorazowej analizy przeznaczenia systemu oraz spełnienia zaawansowanych wymogów przedwdrożeniowych i monitorowania jego działania.
- Kluczowym obszarem staje się zarządzanie ryzykiem, obejmujące jakość danych, przejrzystość działania algorytmów, dokumentację techniczną oraz możliwość audytu podejmowanych decyzji.
- Wprowadzenie obowiązkowego nadzoru człowieka nad systemami AI oznacza konieczność projektowania rozwiązań w taki sposób, aby użytkownik był w stanie zrozumieć działanie systemu, zweryfikować jego wynik i w razie potrzeby podjąć decyzję odmienną.
- Transparentność wobec obywatela staje się standardem funkcjonowania usług publicznych, obejmując obowiązek informowania o wykorzystaniu AI oraz zapewnienia możliwości wyjaśnienia wpływu algorytmu na decyzję administracyjną.
- Wdrożenie systemów AI wiąże się z istotnymi kosztami organizacyjnymi i operacyjnymi, w tym koniecznością budowy kompetencji w zakresie prawa technologicznego, zarządzania danymi oraz bezpieczeństwa informacji.
- Istotnym ryzykiem operacyjnym jest zjawisko tzw. shadow AI, czyli wykorzystywania nieautoryzowanych narzędzi opartych na sztucznej inteligencji, co może prowadzić do naruszeń bezpieczeństwa danych oraz przepisów prawa.
- AI Act funkcjonuje w ścisłym powiązaniu z innymi regulacjami, w szczególności RODO oraz dyrektywą NIS 2, co powoduje, że zarządzanie systemami AI staje się integralną częścią zarządzania bezpieczeństwem informacji i odpowiedzialnością organizacyjną.

- Kluczowym kierunkiem dla administracji publicznej staje się budowa suwerenności technologicznej, rozumianej jako kontrola nad infrastrukturą, danymi oraz wykorzystywanymi modelami, w szczególności poprzez rozwój rozwiązań lokalnych.

Pomimo zwiększonych wymagań regulacyjnych, AI Act stanowi narzędzie budowania zaufania obywateli do państwa, poprzez ograniczenie ryzyk, zwiększenie przejrzystości oraz zapewnienie ochrony praw podstawowych.



Dyrektywa NIS 2 - nowy paradygmat cyberbezpieczeństwa państwa i biznesu

Unijna dyrektywa NIS 2 (Dyrektywa Parlamentu Europejskiego i Rady (UE) 2022/2555) wyznacza nowy, znacznie bardziej rygorystyczny standard w zakresie cyberbezpieczeństwa w Europie. 17 października 2024 roku był terminem, do którego państwa członkowskie zobowiązane były do wdrożenia jej zapisów do krajowych porządków prawnych. W praktyce oznacza to początek nowego etapu regulacyjnego, którego implementacja na poziomie krajowym, w tym w Polsce, jest realizowana poprzez nowelizację przepisów o krajowym systemie cyberbezpieczeństwa.

W Polsce kluczowe znaczenie ma projektowana nowelizacja ustawy o krajowym systemie cyberbezpieczeństwa (KSC), która dostosowuje krajowe regulacje do wymogów NIS 2 i określa szczegółowe obowiązki podmiotów publicznych i prywatnych. To właśnie te przepisy będą stanowiły bezpośrednią podstawę funkcjonowania nowych wymogów w praktyce działania instytucji.

W obliczu rosnącej liczby, skali i wyrafinowania cyberataków, które zagrażają funkcjonowaniu rynków, usług i całych społeczeństw, Unia Europejska podniosła poziom ambicji, przekształcając cyberbezpieczeństwo z obszaru dobrych praktyk w jeden z kluczowych obowiązków regulacyjnych⁰¹.

Czym jest dyrektywa NIS 2?

Dyrektywa NIS 2 wprowadza zaktualizowane ramy prawne, których celem jest ograniczenie rozbieżności w poziomie cyberbezpieczeństwa pomiędzy państwami członkowskimi Unii Europejskiej. Jej założenia muszą zostać wdrożone do krajowych porządków prawnych, dlatego praktyczne znaczenie dla podmiotów publicznych i prywatnych mają przepisy implementujące ją na poziomie krajowym.

Architektura NIS 2 opiera się na nowym podejściu do kwalifikacji podmiotów obję-

tych regulacją. Prawodawca unijny odszedł od wcześniejszego, bardziej uznaniowego modelu identyfikacji operatorów usług kluczowych na rzecz szerszych i bardziej obiektywnych kryteriów, obejmujących wielkość podmiotu oraz sektor, w którym prowadzi on działalność. W efekcie regulacja obejmuje liczne średnie i duże podmioty działające w sektorach uznanych za kluczowe lub ważne z punktu widzenia funkcjonowania państwa i gospodarki.

Podmioty te zostały podzielone na dwie główne kategorie: podmioty kluczowe oraz podmioty ważne. Podział ten wpływa na zakres nadzoru, obowiązków oraz mechanizmów egzekwowania przepisów na poziomie krajowym.

Szanse: Proaktywna ochrona i zabezpieczenie łańcuchów dostaw

Implementacja dyrektywy NIS 2 niesie ze sobą potężny potencjał modernizacyjny, wymuszając transformację podejścia z reaktywnego gaszenia pożarów na proaktywne zarządzanie ryzykiem algorytmicznym i sieciowym. Jedną z najważniejszych szans, jakie kreują nowe przepisy, jest bezprecedensowy nacisk na bezpieczeństwo całego łańcucha dostaw⁰². W minionych latach zaobserwowano, że to właśnie najmniejsze firmy, z racji słabszych zabezpieczeń, stanowią doskonały

wektor ataku dla hakerów dążących do infiltracji infrastruktury krytycznej państw czy wielkich korporacji. Dlatego obecnie podmioty kluczowe i ważne nie mogą skupiać się wyłącznie na własnych serwerach; mają bezwzględny obowiązek weryfikacji i audytowania praktyk cyberbezpieczeństwa swoich dostawców usług, oprogramowania oraz podwykonawców. Wywołuje to „efekt domina”, który systemowo podnosi standardy w całej gospodarce, również u mniejszych kontrahentów⁰³.

Dla rynku europejskiego dyrektywa oznacza także szansę na bezprecedensową standaryzację i wyrównanie szans (tzw. level playing field) dla firm operujących transgranicznie⁰⁴. Ujednolicenie wymogów kończy z dotychczasowym arbitrażem regulacyjnym, w którym firmy mogły rejestrować działalność w krajach o łagodniejszym podejściu do ochrony danych. Ponadto, dokument powołuje do życia nowe, zaawansowane mechanizmy ułatwiające współpracę operacyjną, w tym sieć EU-CyCLONe (Europejska sieć organizacji łącznikowych do spraw kryzysów cyberbezpieczeństwa), mającą koordynować paneuropejskie reagowanie na incydenty na dużą skalę. Zbudowanie odpowiedniej kultury cyberhigieny zminimalizuje zjawisko przestojów operacyjnych i drastycznych strat finansowych, jakie dziś pociągają za sobą skuteczne ataki typu ransomware⁰⁵.

Zagrożenia i wyzwania: Presja czasu, koszty, biurokracja i odpowiedzialność

Transformacja ukształtowana przez NIS 2 pociąga za sobą gigantyczne wyzwania operacyjne, logistyczne oraz finansowe. Kluczowym zagrożeniem, z jakim mierzy się dziś zarówno sektor prywatny, jak i publiczny, jest gwałtowny deficyt wykwalifikowanych specjalistów do spraw cyberbezpieczeństwa. Wdrożenie złożonych systemów zarządzania

bezpieczeństwem informacji (SZBI), w tym mechanizmów uwierzytelniania wieloskładnikowego, segmentacji sieci czy ciągłego monitoringu logów, wymaga ogromnych nakładów finansowych i kapitału ludzkiego.

Największym „batem” regulacyjnym są jednak niezwykle rygorystyczne ramy czasowe raportowania incydentów, które organizacyjnie mogą przerosnąć niejedną firmę. Zainfekowane podmioty mają zaledwie 24 godziny na wyemitowanie wczesnego ostrzeżenia o poważnym zagrożeniu, 72 godziny na przesłanie pełnego zgłoszenia, a następnie jeden miesiąc na przedstawienie kompleksowego raportu końcowego. Brak odpowiednich, zautomatyzowanych procedur w tym zakresie skutkować będzie nałożeniem wysokich kar. Warto doprecyzować, że uchwalona 23 stycznia 2026 r. nowelizacja ustawy o KSC przewiduje istotne przepisy przejściowe w obszarze technologicznym. Obowiązek pełnego korzystania z docelowego systemu teleinformatycznego (S46) do raportowania incydentów i wymiany informacji został odroczony o 12 miesięcy od dnia wejścia w życie przepisów (tj. do 3 kwietnia 2027 r. dla obecnych podmiotów). Daje to jednostkom administracji niezbędny czas na dostosowanie procedur operacyjnych i przeszkolenie kadr, przy jednoczesnym zachowaniu ciągłości raportowania incydentów dotychczasowymi kanałami. System kar finansowych przewidziany w nowych przepisach przypomina te znane z RODO: dla podmiotów kluczowych sankcje mogą sięgnąć 10 milionów euro lub 2% rocznego globalnego obrotu, natomiast dla podmiotów ważnych – 7 milionów euro lub 1,4% obrotu⁰⁶.

Fundamentalnym i najbardziej kontrowersyjnym wyzwaniem jest rewolucja w zakresie odpowiedzialności osobistej kierownictwa. Dyrektywa NIS 2 oraz implementująca ją polska ustawa o krajowym systemie cyberbezpieczeństwa (KSC) wprost wskazują, że za zarządzanie cyberbezpieczeństwem

odpowiada najwyższe kierownictwo, czyli członkowie zarządów, prezydenci miast czy dyrektorzy generalni instytucji⁰⁷.

To oni muszą osobiście zatwierdzać plany zarządzania ryzykiem, a w przypadku rażących zaniedbań ponoszą odpowiedzialność finansową (nowelizacja KSC zakłada kary do 300% miesięcznego wynagrodzenia, a dla podmiotów publicznych do 100%) oraz narażają się na czasowy zakaz pełnienia funkcji zarządczych⁰⁸.

Perspektywa sektora regulowanego i administracji publicznej

Z perspektywy administracji publicznej oraz sektorów regulowanych, nowa dyrektywa stanowi absolutną redefinicję sposobu funkcjonowania. Sektory, które wcześniej rzadko kojarzyły się z wymogami najwyższych standardów IT – z dnia na dzień zostały wpisane na listę podmiotów ważnych lub kluczowych. Szacuje się, że w Polsce nowe przepisy obejmą gwałtownie poszerzoną grupę ok. 40 000 podmiotów⁰⁹¹⁰.

Włączenie do regulacji szeroko pojętej administracji publicznej (zarówno szczebla centralnego, jak i regionalnego czy lokalnego, w tym jednostek samorządu terytorialnego

i spółek komunalnych) nakłada na państwo obowiązek zabezpieczenia danych obywateli na bezprecedensową skalę¹¹. Przerwanie dostaw wody w gminie, paraliż systemu ratownictwa medycznego lub ataki na lokalne sieci transportowe wywołują bowiem natychmiastowy, niszczący efekt kaskadowy. Podmioty te muszą odtąd nie tylko zarządzać incydentami za pośrednictwem krajowego systemu S46 (pamiętając o 12-miesięcznym okresie przejściowym na jego bezwzględne stosowanie), ale także wdrożyć zaawansowane procedury zarządzania ciągłością działania (BCP), testowania scenariuszy odtworzeniowych (DRP) oraz utrzymywania zasobów chmurowych¹².

Dla sektorów silnie regulowanych, wdrażanie NIS 2 wiąże się również z koniecznością prowadzenia obowiązkowych audytów (dla podmiotów kluczowych odbywają się one minimum raz na trzy lata, z pierwszym audytem po 24 miesiącach od wejścia w życie ustawy). Należy pamiętać, że administracja publiczna znajduje się w newralgicznym punkcie również dlatego, że wdrażanie cyberbezpieczeń musi postępować w warunkach ostrego rygoru dyscypliny finansów publicznych, co wymaga mądrego korzystania z unijnych i krajowych funduszy na cyberbezpieczeństwo¹³.

Wnioski podsumowujące

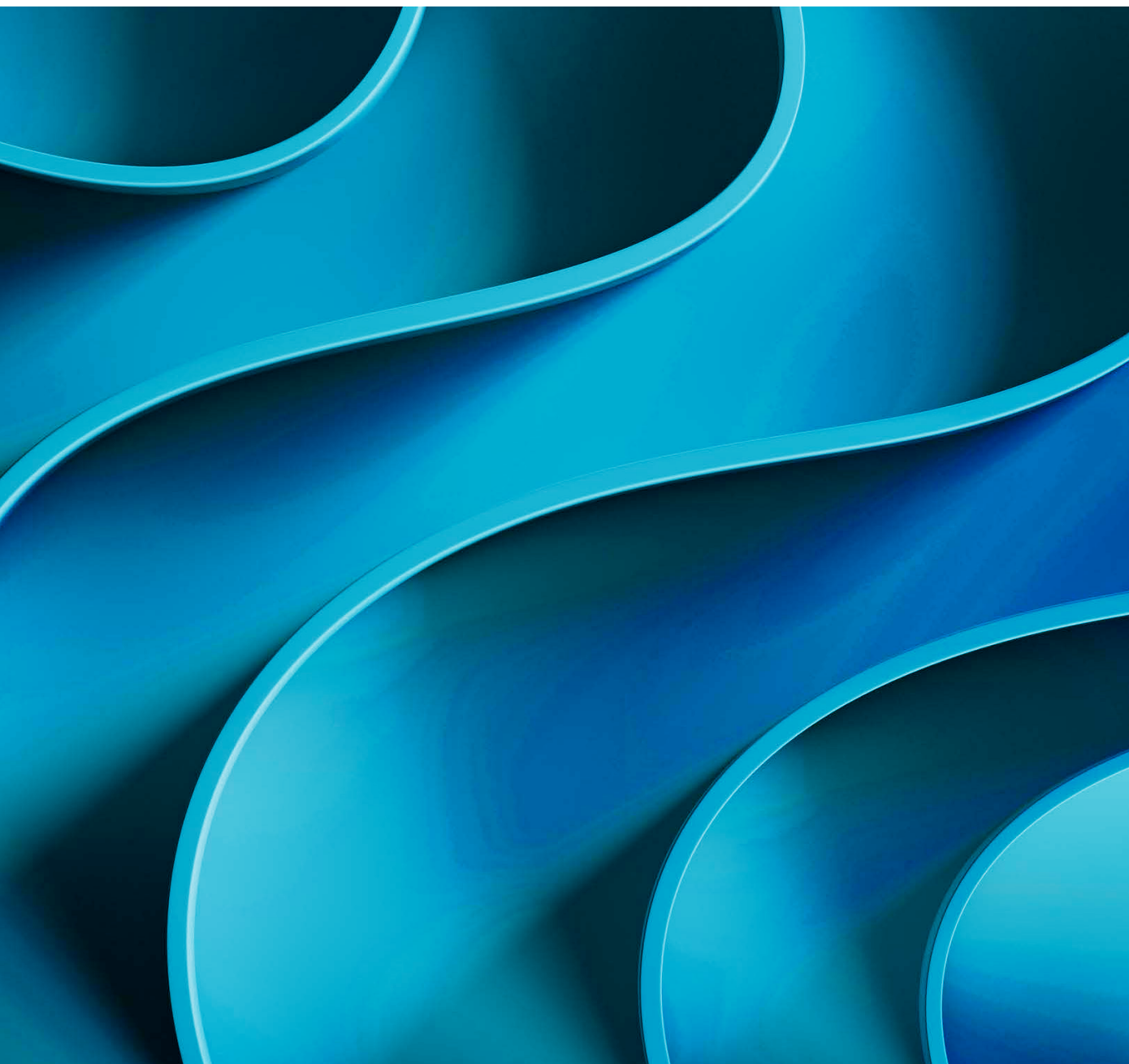
Dyrektywa NIS 2 wprowadza nowy, znacznie bardziej rygorystyczny model zarządzania cyberbezpieczeństwem, który bezpośrednio obejmuje administrację publiczną oraz jednostki samorządu terytorialnego. Regulacja ta redefiniuje podejście do bezpieczeństwa cyfrowego – z obszaru technicznego wsparcia IT staje się ono jednym z kluczowych elementów zarządzania państwem i ciągłością działania usług publicznych.

Z perspektywy administracji publicznej oraz JST kluczowe znaczenie mają następujące wnioski:

- Cyberbezpieczeństwo staje się obowiązkiem systemowym, a nie opcjonalnym elementem infrastruktury IT. Instytucje publiczne muszą traktować je jako integralną część zarządzania organizacją i świadczenia usług publicznych.
- Znacząco rozszerza się zakres podmiotów objętych regulacją. Administracja centralna, jednostki samorządu terytorialnego oraz podmioty komunalne zostają włączone do kategorii podmiotów kluczowych lub ważnych, co oznacza bezpośrednie podleganie rygorystycznym wymagom.
- Następuje przejście z podejścia reaktywnego na proaktywne zarządzanie ryzykiem. Organizacje nie mogą ograniczać się do reagowania na incydenty, lecz muszą aktywnie identyfikować zagrożenia, wdrażać środki zapobiegawcze oraz stale monitorować poziom bezpieczeństwa.
- Kluczowym obszarem staje się bezpieczeństwo łańcucha dostaw. Jednostki publiczne są zobowiązane do weryfikacji dostawców usług, oprogramowania i partnerów technologicznych, co oznacza rozszerzenie odpowiedzialności poza własną infrastrukturę.
- Wprowadzone zostają bardzo restrykcyjne obowiązki raportowania incydentów. Administracja musi być przygotowana do zgłaszania incydentów w krótkich, ściśle określonych terminach (24h, 72h, raport końcowy), co wymaga wdrożenia odpowiednich procedur i automatyzacji.
- Dyrektywa wprowadza bezpośrednią odpowiedzialność kadry zarządzającej. Kierownictwo instytucji publicznych, w tym osoby pełniące funkcje zarządcze w JST, odpowiadają za wdrożenie i nadzór nad systemami cyberbezpieczeństwa, również na poziomie osobistym.
- Wdrożenie NIS 2 wiąże się z istotnymi wyzwaniami organizacyjnymi i finansowymi. Konieczne jest budowanie kompetencji, wdrażanie systemów zarządzania bezpieczeństwem informacji oraz inwestycje w infrastrukturę i zasoby ludzkie, przy jednoczesnym ograniczeniu budżetowym charakterystycznym dla sektora publicznego.
- Niezbędne staje się wdrożenie kompleksowych mechanizmów zapewnienia ciągłości działania (BCP) oraz odtwarzania systemów (DRP), szczególnie w kontekście usług krytycznych takich jak transport, zdrowie czy gospodarka komunalna.

- Administracja publiczna musi przygotować się na regularne audyty oraz stały nadzór regulacyjny, co wymaga uporządkowania procesów, dokumentacji oraz struktury odpowiedzialności w organizacji.
- Wzmacniana jest współpraca na poziomie krajowym i europejskim, w tym poprzez mechanizmy koordynacji reagowania na incydenty, co oznacza większą integrację działań administracji w zakresie cyberbezpieczeństwa.

W konsekwencji dyrektywa NIS 2 wymusza na administracji publicznej i jednostkach samorządu terytorialnego przejście do modelu zarządzania, w którym bezpieczeństwo cyfrowe staje się jednym z fundamentów funkcjonowania państwa. Oznacza to konieczność strategicznego podejścia do cyberbezpieczeństwa, obejmującego zarówno technologie, procesy, jak i odpowiedzialność organizacyjną, przy jednoczesnym zapewnieniu ciągłości i stabilności usług publicznych.



Komentarz eksperta



dr Małgorzata Olszewska

**dr nauk prawnych, autorka publikacji naukowych
w obszarze prawa ochrony danych osobowych
w Internecie.**

Zastępca Dyrektora NASK, Dyrektor Centrum
ds. Projektów Infrastrukturalnych i Edukacyjnych.

Z perspektywy administracji publicznej kwestie regulacyjne oraz odpowiedzialność za wykorzystanie nowych technologii mają dziś charakter fundamentalny. Instytucje publiczne operują na najbardziej wrażliwych danych, dotyczących zdrowia, sytuacji finansowej, rodzinnej czy prawnej obywateli. To właśnie dlatego poziom ochrony informacji w sektorze publicznym musi być szczególnie wysoki.

W ostatnich latach obserwujemy dynamiczny wzrost zagrożeń w cyberprzestrzeni. Ataki ransomware, próby wyłudzeń danych, a także nowe zjawiska, takie jak deepfake czy dezinformacja wspierana przez sztuczną inteligencję, pokazują, że ochrona danych i infrastruktury cyfrowej staje się jednym z kluczowych elementów bezpieczeństwa państwa. W tym kontekście wdrażanie rozwiązań opartych na AI nie może odbywać się w sposób spontaniczny ani niekontrolowany. Jednocześnie należy podkreślić, że regulacje nie powinny być postrzegane jako bariera dla rozwoju technologii. Ich rolą jest stworzenie bezpiecznych i przewidywalnych ram, w których administracja może korzystać z potencjału sztucznej inteligencji. Odpowiednio zaprojektowane przepisy, takie jak AI Act, dyrektywa NIS 2 czy RODO, mają na celu stworzenie ram prawnych bezpiecznego, etycznego i godnego zaufania wdrażania technologii AI, definiując zakres odpowiedzialności oraz wprowadzając mechanizmy nadzoru niezbędne przy pracy z danymi obywateli.

Z punktu widzenia praktyki działania urzędów szczególnie istotne jest zapewnienie równowagi pomiędzy bezpieczeństwem a efektywnością. Administracja nie może rezygnować z wdrażania nowoczesnych narzędzi tylko z obawy przed ryzykiem, ale też nie może wdrażać technologii kosztem utraty kontroli nad danymi. Kluczowe jest takie ułożenie procesów i narzędzi, które umożliwią korzystanie z AI w sposób bezpieczny, transparentny i zgodny z obowiązującymi regulacjami. Z tej perspektywy kluczowym kierunkiem rozwoju jest budowa rozwiązań opartych na zaufanej i kontrolowanej infrastrukturze, w szczególności takich, które umożliwiają przetwarzanie danych w środowiskach lokalnych lub suwerennych. Tylko takie podejście daje administracji realną kontrolę nad informacją i pozwala pogodzić potrzebę innowacji z obowiązkiem ochrony obywatela.

Podsumowując, regulacje w obszarze sztucznej inteligencji są niezbędnym fundamentem jej bezpiecznego wdrażania w sektorze publicznym. Ich celem nie jest ograniczanie rozwoju, lecz umożliwienie jego prowadzenia w sposób odpowiedzialny. Administracja publiczna potrzebuje dziś nie tylko technologii, ale przede wszystkim jasnych zasad, które pozwolą z niej korzystać z pełnym zaufaniem, zarówno ze strony instytucji, jak i obywateli.

RODO w kontekście AI - Ochrona prywatności jako fundament, a nie bariera dla innowacji

Fałszywy mit o barierze dla innowacji

Sztuczna inteligencja z definicji opiera się na przetwarzaniu dużych zbiorów danych, co w przypadku jej wdrażania w sektorze publicznym oznacza bezpośrednie operowanie na informacjach dotyczących obywateli. W debacie technologicznej często pojawia się teza, że restrykcyjne przepisy europejskiego Ogólnego Rozporządzenia o Ochronie Danych (RODO) stanowią barierę dla rozwoju cyfrowego państwa. Aktualne podejście organów nadzorczych wskazuje jednak, że jest to założenie uproszczone. Zgodnie z opiniami Europejskiej Rady Ochrony Danych (EROD), zasady RODO nie blokują innowacji, lecz wyznaczają ramy dla odpowiedzialnego i godnego zaufania wykorzystania danych, również w kontekście rozwoju i wdrażania systemów sztucznej inteligencji⁰¹.

W praktyce oznacza to, że zgodność z przepisami RODO oraz regulacjami wynikającymi z AI Act powinna być traktowana jako element spójnego podejścia do zarządzania ryzykiem technologicznym. Ma to szczególne znaczenie w przypadku systemów klasyfikowanych jako „wysokiego ryzyka”, które mogą wpływać na prawa, obowiązki lub sytuację życiową obywateli. Ochrona danych osobowych pozostaje w tym kontekście jednym z fundamentów polityki cyfrowej państwa, powiązanym z gwarancjami wynikającymi z Karty praw podstawowych Unii Europejskiej oraz krajowych regulacji konstytucyjnych⁰²⁰³.

Warto również odnotować, że na poziomie unijnym trwają prace nad inicjatywami legislacyjnymi mającymi na celu dalsze doprecyzowanie i potencjalne dostosowanie ram prawnych do rozwoju nowych technologii. Jednym z takich projektów jest pakiet określany jako „Digital Omnibus”, który zakłada m.in. możliwe zmiany w zakresie interpretacji danych osobowych, zasad podejmowania decyzji w sposób zautomatyzowany oraz podstaw prawnych przetwarzania danych w kontekście rozwoju sztucznej inteligencji. Na obecnym etapie są to jednak propozycje

legislacyjne, których ostateczny kształt oraz zakres zastosowania pozostają przedmiotem prac instytucji unijnych.

Wyzwanie: Minimalizacja danych i Privacy by Design

Kluczowym wyzwaniem technicznym i prawnym dla administracji publicznej jest pogodzenie potężnego zapotrzebowania modeli językowych na tzw. Big Data z rygorystyczną zasadą minimalizacji danych. Zgodnie z RODO, urzędy mogą zbierać i przetwarzać w systemach AI wyłącznie te informacje, które są absolutnie niezbędne do realizacji określonego, legalnego celu⁰⁴. Wprowadzając innowacje, administracja publiczna ma bezwzględny obowiązek stosowania zasady uwzględniania ochrony danych już w fazie projektowania (data protection by design) oraz domyślnej ochrony danych (data protection by default).

Oznacza to, że architektura systemu algorytmicznego od pierwszej linijki kodu musi minimalizować ryzyko identyfikacji jednostki. Zgodnie jednak z wyraźnym zezwoleniem

samego rozporządzenia AI Act (art. 10 ust. 5), prawo dopuszcza jednak pewne innowacyjne wyjątki: w celu skutecznego wykrywania i korygowania stronniczości algorytmów (tzw. biasu), która mogłaby prowadzić do dyskryminacji, dostawcy systemów mogą wyjątkowo przetwarzać szczególne kategorie danych osobowych⁰⁵. Tego typu operacje obwarowane są jednak potężnymi zabezpieczeniami – dane te muszą być izolowane, pseudonimizowane i usuwane natychmiast po wyeliminowaniu błędu algorytmu.

Tymczasem realia administracji samorządowej wskazują na potężne ryzyko operacyjne. Praktyka pokazuje, że pracownicy urzędów oraz miejskich spółek aktywnie wykorzystują ogólnodostępne czaty AI, takie jak ChatGPT, Copilot czy Gemini, do analizy i streszczania służbowych dokumentów. Przesyłanie danych do takich otwartych narzędzi to prosta droga do naruszeń RODO⁰⁶. Aby temu zapobiec, administratorzy w jednostkach samorządu terytorialnego muszą upewnić się, że wdrażanie AI odbywa się w bezpiecznych, zamkniętych konfiguracjach, z rygorystyczną kontrolą retencji, rejestracją zdarzeń oraz po uprzednim przeprowadzeniu oceny skutków dla ochrony danych (DPIA)⁰⁷.

Uzasadniony interes i ryzyko nielegalnego trenowania modeli

Sztuczna inteligencja, aby osiągnąć wysoką skuteczność, wymaga dostępu do gigantycznych zbiorów informacji. Każda decyzja o zbudowaniu, wdrożeniu lub trenowaniu modelu AI przez polski sektor publiczny czy prywatny musi jednak opierać się na niepodważalnej, solidnej podstawie prawnej. W skomplikowanych technologicznie projektach, w których administracja lub biznes powołuje się na przesłankę „uzasadnionego interesu” jako podstawy do przetwarzania danych osobowych (zgodnie z RODO), konieczne jest uprzednie przeprowadzenie niezwykle drobiazgowego, trzyetapowego testu

bilansującego⁰⁸.

Zgodnie z najnowszymi wytycznymi Europejskiej Rady Ochrony Danych (EROD), test ten ma za zadanie dowieść, że korzyści systemowe płynące z wdrożenia AI są obiektywnie nadrzędne w stosunku do potencjalnych zagrożeń dla podstawowych praw i wolności jednostki⁰⁹. Za uzasadniony interes można uznać na przykład uruchomienie zaawansowanego wirtualnego asystenta (agenta konwersacyjnego) ułatwiającego obywatelom nawigację po usługach publicznych czy implementację algorytmów poprawiających cyberbezpieczeństwo infrastruktury państwowej. Aby jednak takie wdrożenie było legalne, przetwarzanie danych musi zostać uznane za absolutnie niezbędne¹⁰.

Organy nadzoru ochrony danych, w tym polski UODO, będą w takich sytuacjach rygorystycznie badać szereg kryteriów kontekstowych. Ocenie podlegać będzie to, czy dane osobowe wykorzystane do treningu były wcześniej publicznie dostępne, jaki był pierwotny kontekst i źródło ich zebrania, a także czy obywatele mogli racjonalnie oczekiwać i być rzeczywiście świadomi, że ich informacje znajdujące się w internecie posłużą do zasilenia systemów uczących się.

Jeżeli test bilansujący wykaże negatywny wpływ na jednostkę, organizacja ma obowiązek wdrożyć silne środki łagodzące o charakterze technicznym lub organizacyjnym, które zwiększą przejrzystość i ułatwią obywatelom egzekwowanie ich praw¹¹.

DPIA, sankcje finansowe i rola UODO

Dla urzędów i podmiotów publicznych wdrażanie zaawansowanej analityki danych oznacza również ścisłe obowiązki formalne. Przed uruchomieniem technologii konieczne jest przeprowadzenie Oceny Skutków dla Ochrony Danych (DPIA) – jest to bezwzględnie wymagane, gdy operacje obejmują wysokie

ryzyko naruszenia praw i wolności, na przykład przy przetwarzaniu danych medycznych pacjentów lub danych biometrycznych¹². W przypadku naruszenia ochrony danych systemy te wymuszają powiadomienie odpowiednich organów nadzorczych w ciągu 72 godzin¹³. Niezastosowanie się do tych wymogów stwarza podwójne ryzyko: drastycznie podważa zaufanie społeczne do cyfrowego państwa i grozi nałożeniem gigantycznych kar finansowych, które zgodnie z RODO mogą sięgać do 20 milionów euro lub 4% rocznego światowego obrotu podmiotu odpowiedzialnego za wdrażaną technologię. Architektura nadzoru nad tym ekosystemem przewiduje kluczową rolę Urzędu Ochrony Danych Osobowych (UODO). Zgodnie ze zaktualizowaną „Polityką rozwoju sztucznej inteligencji w Polsce do 2030 roku”, UODO będzie sprawował ścisły nadzór nad rynkiem

AI, gwarantując, że cyfrowa transformacja nie odbędzie się kosztem prywatności obywateli. Zapewnienie tych standardów to dziś nie dodatek do technologii, lecz bezwzględny warunek brzegowy jej istnienia w administracji publicznej¹⁴. Dobre praktyki wymagają przy tym traktowania modeli AI jako oddzielnych aktywów informacyjnych w narzędziach do szacowania ryzyka (np. GDPR Risk Tracker) i rygorystycznego mapowania w Rejestrach Czynności Przetwarzania (RCP) specyficznych zagrożeń, takich jak uprzedzenia algorytmiczne (bias), „czarna skrzynka” (brak wyjaśnialności) czy deanonimizacja. Udokumentowana analiza to twarda podstawa wykazania rozliczalności (art. 5 ust. 2 RODO) przy kontrolach UODO.



Wnioski podsumowujące

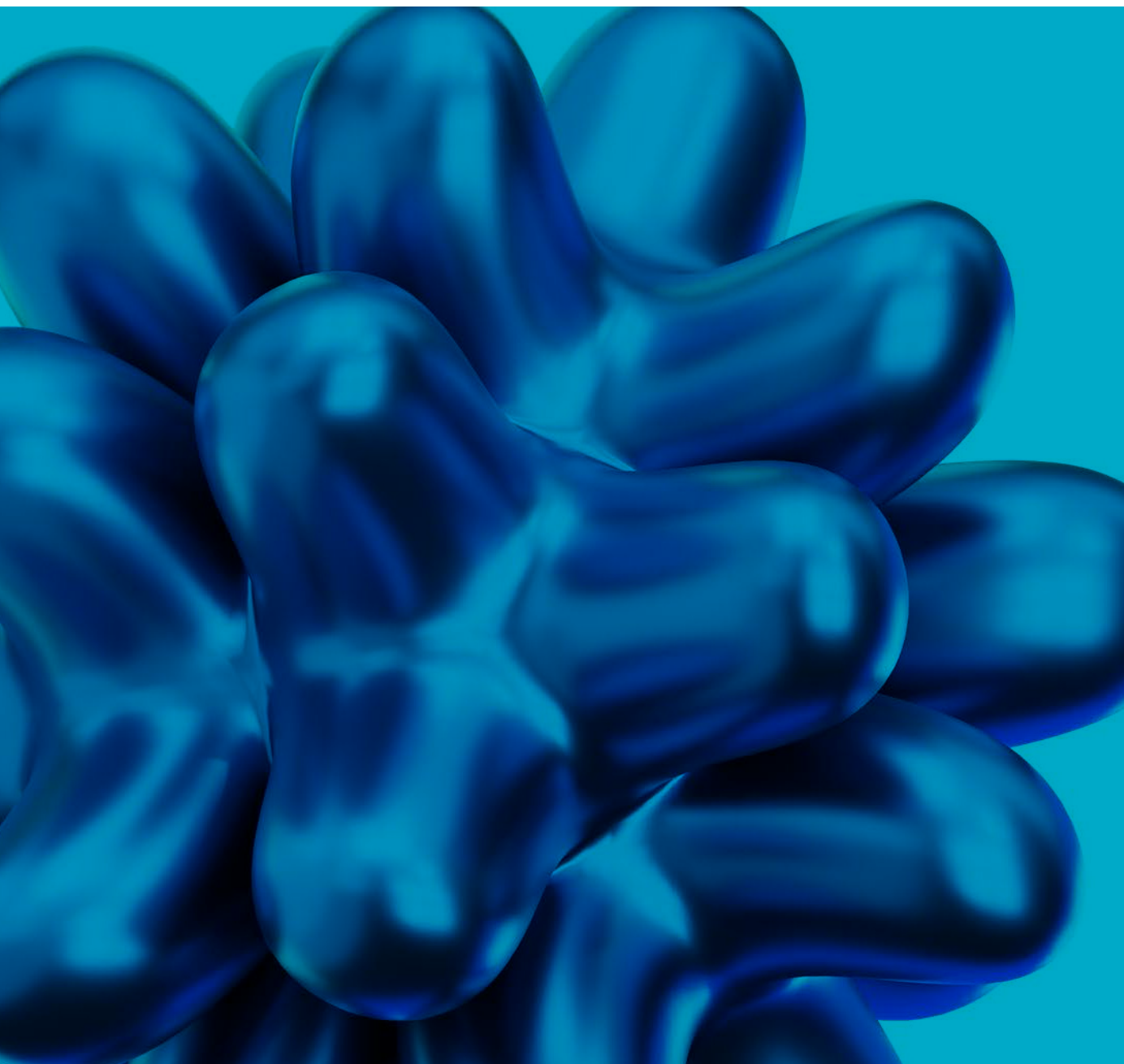
Relacja między rozwojem sztucznej inteligencji a ochroną danych osobowych w sektorze publicznym nie ma charakteru konfliktu, lecz współzależności. Aktualne podejście regulacyjne Unii Europejskiej jasno wskazuje, że zgodność z RODO stanowi fundament bezpiecznego i odpowiedzialnego wdrażania AI, a nie barierę dla innowacji. W praktyce oznacza to konieczność budowy systemów, które od samego początku uwzględniają wymagania prawne oraz ochronę praw obywateli.

Z perspektywy administracji publicznej oraz jednostek samorządu terytorialnego kluczowe znaczenie mają następujące wnioski:

- RODO nie stanowi bariery dla rozwoju AI, lecz tworzy ramy umożliwiające budowę systemów godnych zaufania, które mogą być bezpiecznie wykorzystywane w administracji publicznej.
- Wdrożenia AI w sektorze publicznym muszą być traktowane jako element spójnego systemu regulacyjnego, obejmującego jednocześnie AI Act oraz RODO, szczególnie w kontekście systemów wysokiego ryzyka.
- Kluczowym wyzwaniem operacyjnym jest pogodzenie zapotrzebowania systemów AI na duże zbiory danych z zasadą minimalizacji danych, co wymaga precyzyjnego projektowania procesów przetwarzania informacji.
- Zasada privacy by design oraz privacy by default staje się obowiązkowym standardem, co oznacza konieczność uwzględniania ochrony danych już na etapie projektowania systemów i architektury technologicznej.
- W określonych przypadkach dopuszczalne jest przetwarzanie szczególnych kategorii danych w celu eliminowania błędów i uprzedzeń algorytmicznych, jednak wyłącznie przy zastosowaniu rygorystycznych zabezpieczeń technicznych i organizacyjnych.
- Istotnym ryzykiem operacyjnym w administracji jest wykorzystywanie publicznych narzędzi AI do przetwarzania danych służbowych, co może prowadzić do naruszeń RODO i utraty kontroli nad informacjami.
- Wdrożenia AI w sektorze publicznym wymagają stosowania bezpiecznych, zamkniętych środowisk przetwarzania danych oraz pełnej kontroli nad retencją, dostępem i rejestrowaniem operacji.
- Każde wykorzystanie danych osobowych w systemach AI musi opierać się na jednoznacznej podstawie prawnej, a w przypadku powoływania się na uzasadniony interes konieczne jest przeprowadzenie szczegółowego testu bilansującego.
- Organy nadzorcze będą weryfikować nie tylko podstawę prawną przetwarzania danych, ale również kontekst ich pozyskania, świadomość obywateli oraz realny wpływ na ich prawa i wolności.
- W przypadku wysokiego ryzyka naruszenia praw jednostki obowiązkowe jest przeprowadzenie oceny skutków dla ochrony danych (DPIA), co staje się standardowym elementem wdrażania systemów AI w administracji.

- Naruszenia przepisów RODO wiążą się z istotnym ryzykiem finansowym oraz reputacyjnym, a brak zgodności może podważyć zaufanie obywateli do instytucji publicznych.
- W obszarze ochrony danych osobowych kluczową rolę nadzorczą zachowuje Urząd Ochrony Danych Osobowych. Jednocześnie nadzór nad wdrażaniem i stosowaniem AI nie będzie ograniczał się wyłącznie do UODO, lecz będzie wynikał również z reżimu AI Act oraz właściwości innych organów krajowych i unijnych.

W konsekwencji wdrażanie sztucznej inteligencji w administracji publicznej musi odbywać się w ścisłym powiązaniu z zasadami ochrony danych osobowych, które stanowią nie tylko wymóg prawny, lecz również fundament budowania zaufania do państwa cyfrowego. Oznacza to konieczność projektowania rozwiązań AI w sposób odpowiedzialny, transparentny i zgodny z interesem obywatela, przy jednoczesnym zachowaniu wysokiej efektywności operacyjnej.



KSC i krajowe regulacje - Nowa architektura cyberbezpieczeństwa państwa

Fundamentalna redefinicja bezpieczeństwa cyfrowego

Wdrażanie zaawansowanych systemów informatycznych oraz modeli sztucznej inteligencji w administracji publicznej nie ma szans powodzenia bez niezwykle solidnego fundamentu, jakim jest absolutnie bezpieczna i odporna infrastruktura teleinformatyczna. Przełomowym momentem dla budowy takiej cyfrowej twierdzy w Polsce stało się sfinalizowanie prac nad nowelizacją ustawy o krajowym systemie cyberbezpieczeństwa (KSC), implementującą do polskiego porządku prawnego unijną dyrektywę NIS 2 (Dyrektywa 2022/2555)⁰¹⁰². Akt ten, po wieloletnim, złożonym i miejscami burzliwym procesie legislacyjnym, został ostatecznie podpisany przez Prezydenta RP 19 lutego 2026 roku, co oficjalnie otwiera nowy rozdział w ochronie państwa⁰³. Zmiana ta zdecydowanie nie stanowi jedynie technicznej aktualizacji dotychczasowych, przestarzałych procedur. Jest to fundamentalna redefinicja całego paradygmatu ochrony zasobów informacyjnych – cyberbezpieczeństwo z wysoce wyspecjalizowanej domeny, postrzeganej dotychczas jako wyłączne zadanie i problem działów IT, staje się od teraz rygorystycznym, strategicznym elementem zarządzania na absolutnie najwyższych szczeblach władzy publicznej i korporacyjnej⁰⁴.

Skala wyzwania: Ekspansja na 40 tysięcy podmiotów i rola samorządów

Najbardziej zauważalnym i brzemiennym w skutki efektem znowelizowanej ustawy KSC jest drastyczne poszerzenie katalogu podmiotów objętych surowymi rygorami cybernetycznymi. Ustawodawca całkowicie odszedł od dotychczasowego, wąskiego podziału na operatorów usług kluczowych (OUK) oraz dostawców usług cyfrowych⁰⁵. Zamiast tego wprowadzono nową, dychotomiczną klasyfikację, dzielącą organizacje na „podmioty kluczowe” oraz „podmioty ważne”. O przypisaniu do konkretnej kategorii decyduje w głównej mierze tak zwana zasada wielkościowa (size-cap rule), oparta na kryteriach dla średnich i dużych przedsiębiorstw. W rezultacie zastosowania tak szerokiego sita regulacyjnego, liczba podmiotów zobowiązanych do ścisłego przestrzegania prawa wzrosła w Polsce z zaledwie kilkuset do około 40 tysięcy różnego rodzaju organizacji publicznych i prywatnych⁰⁶⁰⁷.

Z perspektywy przeciętnego obywatela oraz ciągłości świadczenia usług publicznych, najważniejszą zmianą jest bezpośrednie włączenie w ramy ustawy KSC szeroko pojętej administracji publicznej na szczeblu lokalnym i regionalnym. Nowe przepisy objęły swoim zakresem jednostki samorządu terytorialnego (JST), spółki komunalne, samorządowe zakłady budżetowe, a także placówki ochrony zdrowia, szkoły czy lokalne systemy transportowe⁰⁸. Decyzja ta spotkała się z ożywioną debatą w ramach Komisji Wspólnej Rządu i Samorządu Terytorialnego. W toku tych rozmów ustalono, że konieczne jest zróżnicowanie wymagań stawianych jednostkom administracji samorządowej, tak aby były one adekwatne do ich realnej dojrzałości organizacyjnej i dostępnego potencjału kadrowego. Z tego powodu przepisy przewidują mechanizmy wspólnego wykonywania obowiązków i łączenia zasobów przez jednostki samorządowe⁰⁹.

Samoidentyfikacja, obowiązki operacyjne i System Zarządzania Bezpieczeństwem Informacji

W przeciwieństwie do poprzedniego stanu prawnego, urzędy i przedsiębiorstwa nie mogą już biernie czekać na oficjalną decyzję administracyjną o wpisaniu ich na listę podmiotów podlegających ustawie. Nowelizacja KSC wprowadza obowiązek samoidentyfikacji. Każda organizacja funkcjonująca w sektorach objętych ustawą powinna samodzielnie ocenić swój status według kryteriów ustawowych i złożyć wniosek o wpis do wykazu prowadzonego przez ministra właściwego do spraw informatyzacji.

Uzyskanie statusu podmiotu kluczowego lub ważnego uruchamia lawinę obowiązków operacyjnych. Instytucje te mają 12 miesięcy na rygorystyczne wdrożenie i stałe utrzymywanie Systemu Zarządzania Bezpieczeństwem Informacji (SZBI),. Architektura tego systemu jest niezwykle wymagająca – ustawa narzuca prowadzenie ciągłej, systematycznej analizy ryzyka, zarządzanie podatnościami oraz wdrażanie zaawansowanej kryptografii. Ponadto organizacje muszą wypracować i regularnie testować plany ciągłości działania (Business Continuity Plan – BCP) oraz plany odtworzeniowe (Disaster Recovery Plan – DRP),. Plany te są kluczowe w przypadku ataków typu ransomware, gdyż pozwalają na techniczne odtworzenie zasobów z bezpiecznych kopii zapasowych i utrzymanie świadczenia usług w oparciu o procedury alternatywne¹⁰.

Ochrona łańcucha dostaw i instytucja „Dostawcy Wysokiego Ryzyka” (DWR)

Nowoczesne środowiska informatyczne, ze sztuczną inteligencją na czele, charakteryzują się ogromnym skomplikowaniem i zależno-

ściami od podmiotów trzecich. Z tego powodu znowelizowana ustawa KSC kładzie niespotykany dotąd nacisk na kompleksowe bezpieczeństwo łańcucha dostaw (Supply Chain Security). Podmioty administracji publicznej oraz firmy muszą wnikliwie weryfikować swoich podwykonawców, upewniając się, że dostarczane przez nich usługi chmurowe, oprogramowanie czy modele algorytmiczne nie stanowią luki w zabezpieczeniach¹¹¹².

Aby uszczelnić krajową infrastrukturę krytyczną, polski ustawodawca wyposażył organy państwa w potężne, a zarazem kontrowersyjne narzędzie w postaci procedury „Dostawcy Wysokiego Ryzyka” (DWR),. Zgodnie z nowymi przepisami, Minister Cyfryzacji, w drodze decyzji administracyjnej, może uznać konkretnego dostawcę sprzętu lub oprogramowania za poważne zagrożenie dla bezpieczeństwa Rzeczypospolitej. Kwalifikacja ta opiera się na niejawnych przesłankach technicznych oraz wnikliwej ocenie uwarunkowań geopolitycznych. Skutki takiej decyzji są natychmiastowe i bezwzględne: podmioty kluczowe i ważne zostają prawnie zmuszone do całkowitego usunięcia i wycofania z użytku produktów takiego dostawcy w maksymalnym terminie do 7 lat (lub zaledwie 4 lat dla kluczowych funkcji telekomunikacyjnych)¹³. Co wywołuje najgorętsze dyskusje na rynku biznesowym, przepisy nie przewidują wypłaty żadnych odszkodowań ze Skarbu Państwa za konieczność tak kosztownej i przymusowej wymiany infrastruktury¹⁴. Z uwagi na ostre zarzuty dotyczące możliwego naruszenia swobody prowadzenia działalności gospodarczej oraz znamion „wyłączenia”, mechanizm DWR został skierowany w trybie kontroli następcej do Trybunału Konstytucyjnego. Co więcej, wprowadzono sprzężenie tych przepisów z Prawem Zamówień Publicznych – znowelizowane prawo z urzędu nakazuje odrzucanie ofert przetargowych na sprzęt i usługi ICT pochodzące od dostawców z listy DWR¹⁵.

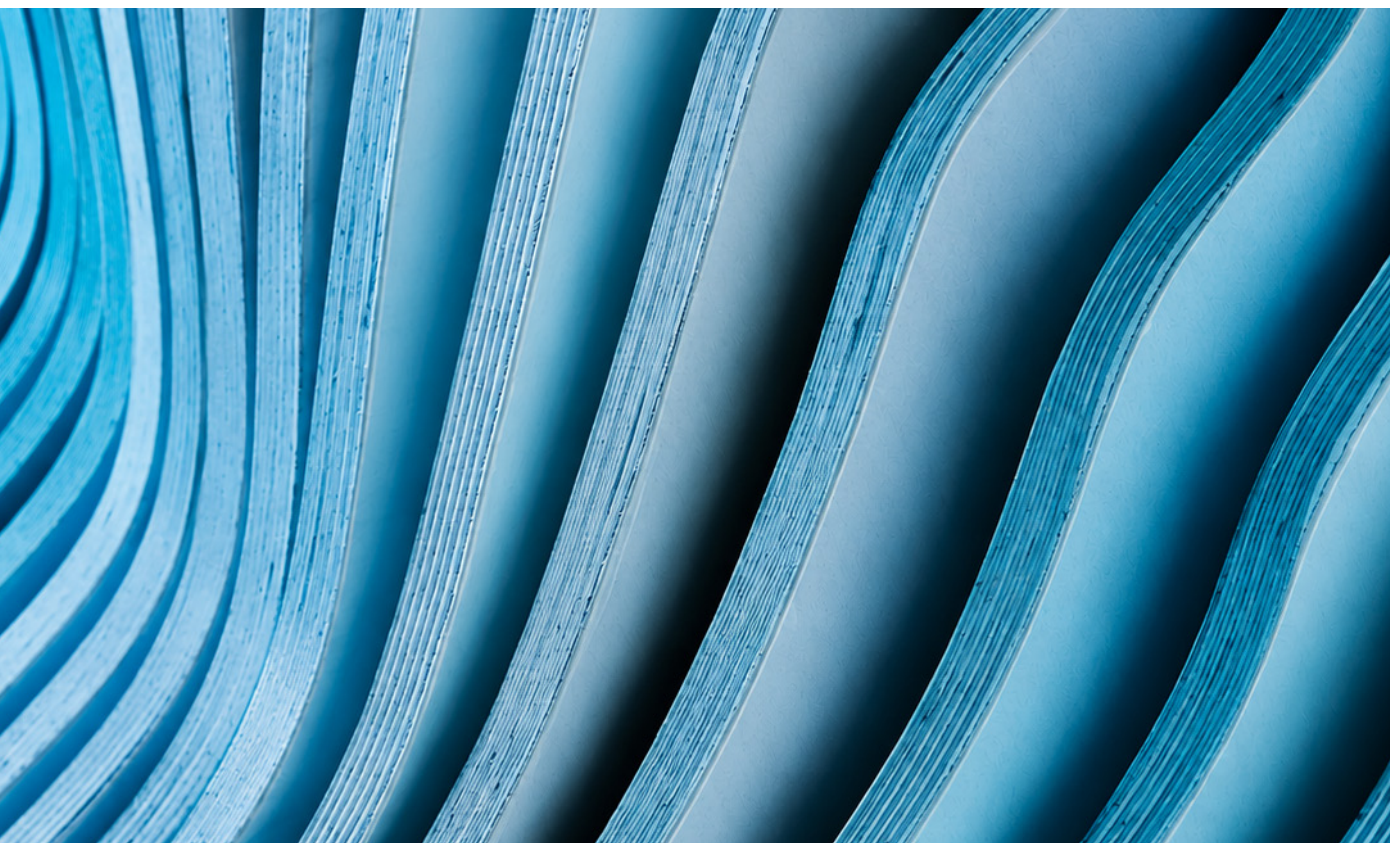
Osobista odpowiedzialność, obligatoryjne audyty i kary finansowe

Aby nowa architektura cyberbezpieczeństwa państwa nie stała się wyłącznie zbiorem teoretycznych i martwych zaleceń, ustawa KSC drastycznie redefiniuje i zaostrza system sankcyjny. Po pierwsze, weryfikacja zgodności z prawem przestała być dobrowolna – każdy podmiot kluczowy ma bezwzględny obowiązek przeprowadzenia głębokiego, niezależnego audytu bezpieczeństwa już w ciągu pierwszych 24 miesięcy od objęcia go ustawą¹⁶.

Po drugie, kary finansowe za niedopełnienie obowiązków osiągnęły poziom zbliżony do sankcji znanych z RODO. W przypadku podmiotów kluczowych, kary nakładane przez organy nadzorcze mogą sięgnąć astronomicznej kwoty 10 milionów euro lub 2% rocznego, globalnego obrotu przedsiębiorstwa. Dla podmiotów ważnych ustawodawca przewidział maksymalne sankcje rzędu 7 milionów euro lub 1,4% przychodów. Ponadto nowelizacja polska narzuca karę nawet do 100 mln PLN za naruszenia skutkujące skrajnie poważnym zagrożeniem (dla obronności,

bezpieczeństwa państwa, życia lub zdrowia ludzi). Organy kontrolne zostały wyposażone w szerokie uprawnienia weryfikacyjne, włączając w to m.in. prawo do żądania na koszt kontrolowanego profesjonalnych tłumaczeń zagranicznej dokumentacji technicznej, co znacząco zwiększa rygor kontroli¹⁷.

Prawdziwym wstrząsem dla sektora publicznego i prywatnego jest jednak wprowadzenie zasady osobistej odpowiedzialności kadry kierowniczej. Od teraz, wójtowie, prezydenci miast, dyrektorzy urzędów oraz członkowie zarządów spółek osobiście odpowiadają za stan zabezpieczeń. Za rażące zaniedbania (np. brak wdrożenia SZBI lub zignorowanie poleceń organów) mogą oni zostać ukarani z prywatnego majątku karą w wysokości do 300% ich miesięcznego wynagrodzenia. W odniesieniu do kierowników samorządowych podmiotów publicznych, kara ta wynosi do 100% miesięcznej pensji, a środki z tych kar zasilają specjalny, państwowy Fundusz Cyberbezpieczeństwa. Dodatkowo, w przypadku ekstremalnych naruszeń, organ nadzorczy może orzec czasowy zakaz pełnienia funkcji zarządczych wobec osób kierujących danym podmiotem¹⁸.



Prawo Zamówień Publicznych (PZP) - Uszczelnianie łańcucha dostaw technologii

Cyfrowa suwerenność i cyberbezpieczeństwo aparatu państwowego nie mogą być realizowane wyłącznie na poziomie deklaracji politycznych. Realnym narzędziem egzekwowania tych standardów są publiczne pieniądze. Wraz z wejściem w życie znowelizowanej ustawy o krajowym systemie cyberbezpieczeństwa (3 kwietnia 2026 r.), uruchomiono również niezwykle rygorystyczne mechanizmy w ustawie z dnia 11 września 2019 r. – Prawo zamówień publicznych (PZP). Zmiany te diametralnie przekształcają procedury przetargowe w administracji, izolując publiczną infrastrukturę, w tym wdrażane systemy sztucznej inteligencji, od niebezpiecznych technologii¹⁹.

Nowa architektura zamówień publicznych wprowadza do ustawy PZP dwie potężne podstawy do obligatoryjnego odrzucania ofert wykonawców w postępowaniach przetargowych. Zgodnie z nowym brzmieniem art. 226 ust. 1 pkt 17 ustawy PZP, zamawiający (urząd, szpital, ministerstwo) ma bezwzględny obowiązek odrzucenia oferty, jeżeli obejmuje ona produkt ICT, usługę ICT lub proces ICT, które zostały wskazane w oficjalnej rekomendacji Pełnomocnika Rządu ds. Cyberbezpieczeństwa jako wywierające negatywny wpływ na podstawowy interes bezpieczeństwa państwa²⁰. Przepis ten w sposób elastyczny pozwala państwu na szybkie reagowanie na ujawniane podatności sprzętowe i programowe, odcinając zakażone lub celowo osłabione systemy od publicznego finansowania.

Jeszcze dalej idącą zmianą, wywołującą potężne skutki dla całego rynku technologicznego w Polsce, jest dodanie do PZP zupełnie nowego art. 226 ust. 1 pkt 19. Przepis ten nakłada na każdego zamawiającego z sektora publicznego bezwzględny obowiązek odrzucenia oferty, jeśli zaproponowane w niej produkty, usługi lub procesy teleinformatyczne pochodzą od podmiotu oficjalnie uznanego za „dostawcę wysokiego ryzyka” (DWR/HRV). Taka kwalifikacja następuje na mocy decyzji administracyjnej Ministra Cyfryzacji, ogłaszanej publicznie w „Monitorze Polskim”.

Należy podkreślić niezwykle restrykcyjny, retrospektywny charakter tych przepisów prawnych. Ustawa nowelizująca nie przewiduje w tym zakresie tzw. przepisów przejściowych. Oznacza to, że nowa podstawa odrzucania ofert ze względu na obecność sprzętu od dostawcy wysokiego ryzyka (art. 226 ust. 1 pkt 19 PZP) ma zastosowanie z urzędu również do wszystkich postępowań przetargowych, które zostały wszczęte i nie zostały zakończone (rozstrzygnięte) przed dniem 3 kwietnia 2026 roku. Urzędnicy prowadzący postępowania muszą więc dokonać ponownej weryfikacji ofert pod kątem nowej, geopolitycznej i technologicznej mapy ryzyka.

Przepisy te wprost realizują założenia zaktualizowanej Strategii Cyberbezpieczeństwa RP, która wymusza uwzględnianie w zamówieniach publicznych rygorów certyfikacji cyberbezpieczeństwa (np. systemów opartych o wspólne kryteria EUCC), stosowania bezpiecznych standardów kryptograficznych oraz – co kluczowe w dobie dążenia do suwerenności – wspierania produktów z zakresu cyberbezpieczeństwa opartych na otwartym oprogramowaniu (open source) oraz krajowych, zaufanych dostawcach²¹. Prawo zamówień publicznych stało się tym samym główną służą bezpieczeństwa, przez którą do polskich urzędów mogą trafić wyłącznie takie modele sztucznej inteligencji i infrastruktura, które gwarantują państwu pełną suwerenność i niezależność od globalnych dostawców.

Wnioski regulacyjne dla administracji publicznej i samorządów

Przełom lat 2024–2026 ukształtował zupełnie nową, wymagającą architekturę prawną dla innowacyjnego państwa. Kluczowe wnioski z analizy przepisów kształtujących ten ekosystem (AI Act, NIS 2, KSC, RODO, PZP) to:

- **Podejście oparte na ryzyku (Risk-Based Approach):** Europejski AI Act oraz dyrektywa NIS 2 wymuszają odejście od uniwersalnych rozwiązań na rzecz analizy punktowej. Istotna część algorytmów stosowanych w usługach publicznych (np. w obszarze świadczeń, rekrutacji, edukacji, ochrony zdrowia czy bezpieczeństwa) może zostać zakwalifikowana jako systemy wysokiego ryzyka, co wymaga każdorazowej analizy przeznaczenia systemu, oceny skutków dla praw podstawowych tam, gdzie jest to wymagane, oraz zapewnienia realnego nadzoru człowieka.
- **Kierownictwo odpowiada osobiście:** Nowelizacja ustawy o Krajowym Systemie Cyberbezpieczeństwa (KSC) kończy z traktowaniem cyberbezpieczeństwa jako wyłącznego problemu działów IT. Kierownictwo podmiotów kluczowych i ważnych (wójtowie, prezydenci miast, dyrektorzy, zarządy) ponosi osobistą odpowiedzialność za wdrożenie adekwatnych systemów zarządzania bezpieczeństwem informacji, a sankcje mogą obejmować również odpowiedzialność finansową osób zarządzających.
- **Ochrona prywatności jako fundament:** RODO, wbrew powszechnym mitom, nie hamuje rozwoju AI. Wymusza jednak zasadę privacy by design, minimalizację danych w zbiorach treningowych oraz bezwzględne obowiązki informacyjne wobec obywateli. Nielegalne użycie danych (np. nieautoryzowany web scraping) może zdyskwalifikować użyteczność całego modelu algorytmicznego.
- **Bezpieczeństwo łańcucha dostaw:** Radykalne zmiany w Prawie Zamówień Publicznych (PZP) i ustawie KSC wprowadzają bezwzględny mechanizm odrzucania technologii od „dostawców wysokiego ryzyka” (DWR). Urzędy zyskują twarde podstawy do wykluczania ofert zagrażających suwerenności informacyjnej kraju i są zmuszone prawnie do weryfikacji cyberbezpieczeństwa na każdym etapie u swoich podwykonawców.
- **Suwerenność technologiczna:** Biorąc pod uwagę gęszcz wymogów, jednym z preferowanych kierunków rozwoju dla administracji jest oparcie transformacji na rozwiązaniach zapewniających realną kontrolę nad infrastrukturą, danymi i modelem przetwarzania – w tym na technologiach lokalnych lub innych rozwiązaniach spełniających równoważne wymogi bezpieczeństwa i zgodności.
- **Aktywna samoocena:** Poszerzenie katalogu KSC obejmuje dziś JST i dziesiątki tysięcy nowych podmiotów. Prawo wymusza proaktywne podejście: samoidentyfikację, złożenie wniosku o wpis do wykazu co do zasady w terminie 6 miesięcy od spełnienia przesłanek uznania za podmiot kluczowy lub ważny oraz przygotowanie do podłączenia do systemu S46 i wdrożenia obowiązków ustawowych w terminach wynikających z ustawy.

Rozdział 3

Regulacje i odpowiedzialność

1. KPMG, Sztuczna inteligencja w Polsce, <https://kpmg.com/pl/pl/wiedza/technologia/sztuczna-inteligencja-w-polsce.html>, dostęp: 25.03.2026.

Akt w sprawie sztucznej inteligencji

1. Parlament Europejski, Akt w sprawie sztucznej inteligencji – pierwsze przepisy regulujące AI, <https://www.europarl.europa.eu/topics/pl/article/20230601STO93804/akt-ws-sztucznej-inteligencji-pierwsze-przepisy-regulujace-ai>, dostęp: 25.03.2026.

2. Parlament Europejski i Rada (UE), Rozporządzenie (UE) 2024/1689 (AI Act), <https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=CELEX:32024R1689>, dostęp: 25.03.2026.

3. Kancelaria Sawaryn i Partnerzy, Systemy wysokiego ryzyka według AI Act, <https://sawaryn.com/publikacje/systemy-wysokiego-ryzyka-wedlug-ai-act/>, dostęp: 25.03.2026.

4. Parlament Europejski, EU AI Act – first regulation on artificial intelligence, <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>, dostęp: 25.03.2026.

5. NASK, Ilu urzędników administracji publicznej korzysta z AI – wyniki raportu, <https://www.prawo.pl/samorzad/ilu-urzednikow-administracji-publicznej-korzysta-z-ai-raport-nask,1540747.html>, dostęp: 25.03.2026

Dyrektywa NIS 2 i nowy paradygmat cyberbezpieczeństwa państwa i biznesu

1. Komisja Europejska, NIS 2 Directive – cybersecurity in the EU, <https://digital-strategy.ec.europa.eu/pl/policies/nis2-directive>, dostęp: 25.03.2026.

2. M. Kaczorowska, Dyrektywa NIS 2 jako narzędzie podnoszenia poziomu cyberbezpieczeństwa w UE, Journal of Modern Science, <https://www.journaldot.pl/pdf-215365-132995?filename=Dyrektywa-NIS-2-jako-narz.pdf>, dostęp: 25.03.2026.

3. Trecom, Dyrektywa NIS 2 w 2026 roku – co musisz wiedzieć, kogo obowiązuje i jakie nakłada obowiązki, <https://www.trecom.pl/dyrektywa-nis-2-w-2026-roku-co-musisz-wiedziec-kogo-obowiazuje-i-jakie-naklada-obowiazki/>, dostęp: 25.03.2026.

4. Parlament Europejski i Rada (UE), Dyrektywa (UE) 2022/2555 (NIS 2), <https://eur-lex.europa.eu/legal-content/PL/ALL/?uri=CELEX:32022L2555>, dostęp: 25.03.2026.

5. Trecom, Dyrektywa NIS 2 – zakres obowiązków i podmioty objęte regulacją, <https://www.trecom.pl/dyrektywa-nis-2-w-2026-roku-co-musisz-wiedziec-kogo-obowiazuje-i-jakie-naklada-obowiazki/>, dostęp: 25.03.2026.

6. EY, Dyrektywa NIS 2 – zakres regulacji i obowiązki organizacji, https://www.ey.com/pl_pl/dyrektywa-nis2, dostęp: 25.03.2026.

7. Ministerstwo Cyfryzacji, Bezpieczniejsze sieci i stabilne usługi – nowelizacja ustawy o krajowym systemie cyberbezpieczeństwa coraz bliżej, <https://www.gov.pl/web/cyfryzacja/bezpieczniejsze-sieci-i-stabilne-uslugi-nowelizacja-ustawy-o-krajowym-systemie-cyberbezpieczenstwa-coraz-blizej>, dostęp: 25.03.2026.

8. NFLO, Które sektory są objęte dyrektywą NIS 2 – przegląd zakresu cyberbezpieczeństwa w UE, <https://nflo.pl/baza-wiedzy/ktore-sektory-sa-objete-dyrektywa-nis2-kompleksowy-przeglad-rozszerzonego-zakresu-cyberbezpieczenstwa-w-ue/>, dostęp: 25.03.2026.

9. BrandsIT, Kary za brak zgodności z NIS2 – jak przygotować firmę na nowe wymogi KSC, <https://brandsit.>

Rozdział 3

pl/kary-za-brak-zgodnosci-z-nis2-jak-przygotowac-firme-na-nowe-wymogi-ksc/, dostęp: 25.03.2026.

10. DSO Consulting, Nowelizacja KSC i NIS 2 w JST – analiza wdrożenia i obowiązków, <https://dsoconsulting.pl/nowelizacja-ksc-nis2-jst-analiza/>, dostęp: 25.03.2026.

11. EY, Dyrektywa NIS 2 – nowe obowiązki w zakresie cyberbezpieczeństwa, https://www.ey.com/pl_pl/dyrektywa-nis2, dostęp: 25.03.2026.

RODO w kontekście AI

1. Europejska Rada Ochrony Danych (EDPB), Opinia dotycząca modeli AI i zasad RODO wspierających odpowiedzialne wykorzystanie sztucznej inteligencji, https://www.edpb.europa.eu/news/news/2024/edpb-opinion-ai-models-gdpr-principles-support-responsible-ai_pl, dostęp: 25.03.2026.

2. Wolters Kluwer, RODO a AI – zgodność i wyzwania regulacyjne (Digital Omnibus), <https://www.wolterskluwer.com/pl-pl/expert-insights/rodo-ai-zgodnosc-digital-omnibus-przewodnik>, dostęp: 25.03.2026.

3. Polski Instytut Ekonomiczny, Sztuczna inteligencja w administracji publicznej, <https://pie.net.pl/wp-content/uploads/2024/09/Sztuczna-inteligencja-w-administracji-publicznej.pdf>, dostęp: 25.03.2026.

4. Kancelaria Sawaryn i Partnerzy, Systemy wysokiego ryzyka według AI Act, <https://sawaryn.com/publikacje/systemy-wysokiego-ryzyka-wedlug-ai-act/>, dostęp: 25.03.2026.

5. Parlament Europejski i Rada (UE), Rozporządzenie (UE) 2024/1689 z dnia 13 czerwca 2024 r. ustanawiające zharmonizowane przepisy dotyczące sztucznej inteligencji (AI Act), <https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=CELEX:32024R1689>, dostęp: 25.03.2026.

6. Instytut Spraw Obywatelskich, Czy AI zastąpi urzędnika? Transformacja, która wymaga kontroli, <https://instytutsprawobywatelskich.pl/czy-ai-zastapi-urzednika-transformacja-ktora-wymaga-kontroli/>, dostęp: 25.03.2026.

7. Ministerstwo Cyfryzacji, Polityka rozwoju sztucznej inteligencji w Polsce do 2030 roku – bezpieczne i odpowiedzialne wykorzystanie AI, <https://www.gov.pl/web/cyfryzacja/polityka-rozwoju-sztucznej-inteligencji-w-polsce-do-2030-roku-bezpieczne-i-odpowiedzialne-wykorzystanie-sztucznej-inteligencji>, dostęp: 25.03.2026.

KSC i krajowe regulacje

1. Parlament Europejski i Rada (UE), Dyrektywa (UE) 2022/2555 (NIS 2), Dz.U. UE L 333 z 27.12.2022, <https://eur-lex.europa.eu/>, dostęp: 25.03.2026.

2. Traple Konarski Podrecki i Wspólnicy, Dyrektywa NIS 2 – podmioty kluczowe i ważne oraz zakres regulacji, <https://www.traple.pl/dyrektywa-nis-2-kogo-obejma-nowe-przepisy-podmioty-kluczowe-i-wazne/>, dostęp: 25.03.2026.

3. CyCommSec, Dyrektywa NIS 2 – wymagania i zakres regulacji cyberbezpieczeństwa, <https://cycommsec.com/nis2>, dostęp: 25.03.2026.

4. Sejm RP, Ustawa o krajowym systemie cyberbezpieczeństwa (KSC), <https://isap.sejm.gov.pl/isap.nsf/DocDetails.xsp?id=WDU20260000252>, dostęp: 25.03.2026.

5. NASK, Obowiązki podmiotów kluczowych i ważnych w dyrektywie NIS 2, <https://cyberpolicy.nask.pl/obowiazki-podmiotow-kluczowych-i-waznych-w-dyrektywie-nis2/>, dostęp: 25.03.2026.

6. M. Kaczorowska, Dyrektywa NIS 2 jako narzędzie podnoszenia poziomu cyberbezpieczeństwa w UE, „Journal of Modern Science”, <https://www.journaldot.pl/pdf-215365-132995?filename=Dyrektywa-NIS-2-jako-narz.pdf>, dostęp: 25.03.2026.

Rozdział 3

7. Supremo, Dyrektywa NIS 2 – nowe obowiązki dla firm, <https://www.supremo.pl/blog-posty/dyrektywa-nis-2-nowe-obowiazki-dla-firm>, dostęp: 25.03.2026.
8. Gazeta Prawna, Nowe przepisy o cyberbezpieczeństwie – konsekwencje regulacji NIS 2, <https://edgp.gazetaprawna.pl/prawo/prawo-internetu-i-ochrony-danych/artykuly/10615291,nowe-przepisy-o-cyberbezpieczenstwie-to-de-facto-wywlaszczenie-bez-ods.html>, dostęp: 25.03.2026.
9. Urząd Zamówień Publicznych, Zmiany w PZP w związku z nowelizacją ustawy o krajowym systemie cyberbezpieczeństwa, <https://www.gov.pl/web/uzp/zmiany-w-pzp-w-zwiazku-z-ogloszeniem-ustawy-o-zmianie-ustawy-o-krajowym-systemie-cyberbezpieczenstwa>, dostęp: 25.03.2026.

Prawo Zamówień Publicznych (PZP)

1. Sejm RP, Ustawa o krajowym systemie cyberbezpieczeństwa (KSC), <https://isap.sejm.gov.pl/isap.nsf/DocDetails.xsp?id=WDU20260000252>, dostęp: 25.03.2026.
2. Urząd Zamówień Publicznych, Zmiany w PZP w związku z nowelizacją ustawy o krajowym systemie cyberbezpieczeństwa, <https://www.gov.pl/web/uzp/zmiany-w-pzp-w-zwiazku-z-ogloszeniem-ustawy-o-zmianie-ustawy-o-krajowym-systemie-cyberbezpieczenstwa>, dostęp: 25.03.2026.
3. Rada Ministrów, Strategia cyberbezpieczeństwa Rzeczypospolitej Polskiej, <https://sip.lex.pl/akty-prawne/mp-monitor-polski/strategia-cyberbezpieczenstwa-rzeczypospolitej-polskiej-22248638>, dostęp: 25.03.2026.

04

AI w administracji publicznej



krajowa
izba
gospodarki
cyfrowej

AI w administracji publicznej - przegląd

Sektor publiczny to fundament funkcjonowania każdego rozwijającego się państwa. Z natury rzeczy, jest to środowisko o bardzo złożonym i trudnym charakterze, obwarowane wielostopniowymi procedurami, przepisami działającymi pod ogromną presją odpowiedzialności społecznej. Każdego dnia administracja ma do czynienia z ogromnymi wolumenami informacji o obywatelach, co sprawia, że każda zmiana technologiczna wymaga rozważań, zachowywania najwyższych standardów bezpieczeństwa i działania na gruncie przepisów i wykładni prawa. Kluczem przed podjęciem jakichkolwiek działań wdrożeniowych jest więc pełne zrozumienie i wypracowanie przejrzystych ram prawnych, które zapewnią urzędnikom pewność oraz poczucie bezpieczeństwa przy korzystaniu z nowych narzędzi. Zaufanie społeczne do cyfrowej transformacji buduje się bowiem poprzez klarowność zasad i jasne przypisanie odpowiedzialności, co stanowi bezpośrednią gwarancję dla obywatela. Tylko ściśle i synergiczne stosowanie obowiązujących regulacji zapewnia pełną przejrzystość procesów i daje obu stronom absolutną pewność, że gromadzone dane są chronione, a fundamentalne prawa mieszkańców pozostają nienaruszone.

Choć współczesna debata publiczna zdominowana jest przez wszechobecne chatboty i generatory obrazów, sztuczna inteligencja wcale nie jest wynalazkiem ostatnich lat. Jej korzenie sięgają głęboko w przeszłość – fundamenty tej dziedziny kładziono już w latach 50. XX wieku, chociażby podczas słynnej konferencji w Dartmouth w 1956 roku, która oficjalnie zapoczątkowała AI jako dziedzinę nauki. To tam właśnie po raz pierwszy sformułowano termin „sztuczna inteligencja” i uznano ją za nową dyscyplinę naukową. Przez kolejne dekady technologia ta rozwijała się w zaciszu laboratoriów badawczych, przechodząc przez okresy następującego po sobie dynamicznego rozwoju i stagnacji. Była to domena wąskiego grona naukowców, programistów i potężnych korporacji technologicznych.

Prawdziwym punktem zwrotnym, który wywołał globalny szok na rynku, była to dopiero końcówka 2022 roku i publiczna premiera ChatGPT, wdrożonego przez firmę OpenAI. Uznaje się, że to właśnie ten moment rozpoczął prawdziwą „rewolucję AI”, udostępniając rozwiązania sztucznej inteligencji w ręce każdego człowieka, nadając tempo rozwoju

i zmian jakiegoś świata wcześniej nie widział.

Jak zauważają Christof Ebert i Panos Louridas w swojej pracy badawczej, na przestrzeni ostatniego stulecia czas potrzebny na masową adaptację przełomowych innowacji drastycznie się kurczy. Aby zdobyć pierwsze 100 milionów użytkowników, tradycyjny telefon potrzebował aż 900 miesięcy, samochód – 400 miesięcy, a telefonia komórkowa – 190 miesięcy. Nawet symbole nowoczesnej cyfryzacji, takie jak globalna sieć internet (80 miesięcy), komunikator WhatsApp (40 miesięcy) czy bijąca rekordy popularności platforma TikTok (9 miesięcy), wydają się powolne w zderzeniu z dynamiką generatywnej sztucznej inteligencji.

Od momentu udostępnienia aplikacji ChatGPT szerokiej publiczności pod koniec 2022 roku, narzędzie to osiągnęło pułap 100 milionów użytkowników w zaledwie 2 miesiące, rozpoczynając tym samym globalny wyścig zbrojeń w dziedzinie AI na niespotykaną wcześniej skalę. Liczby te pokazują, że rozwój sztucznej inteligencji, to nie tylko teoria, to realna zmiana rzeczywistości milionów ludzi na całym świecie. Wobec tych twardych

Czas osiągnięcia 100 mln użytkowników



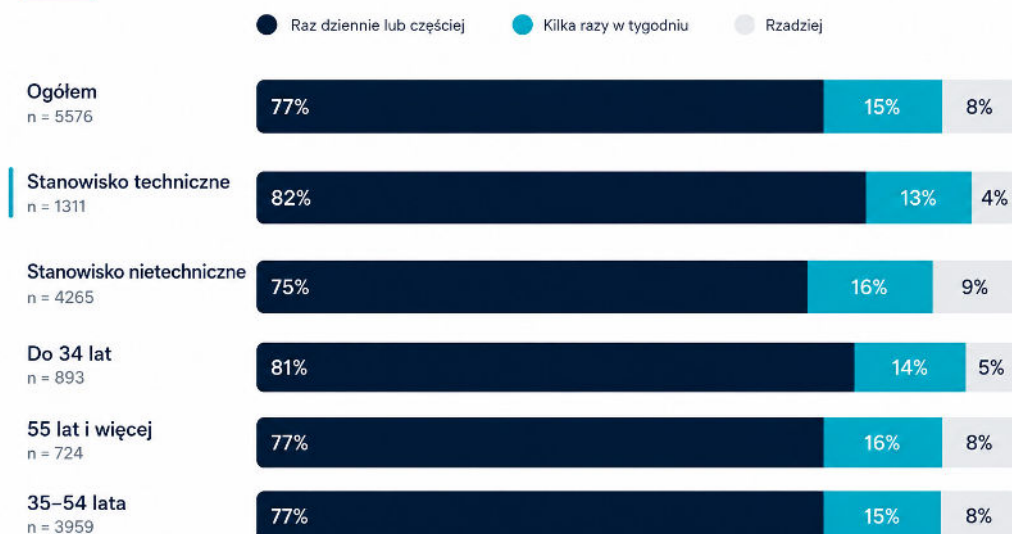
Wykres 4.1 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

danych, nie sztuką jest przejść wokół nich obojętnie. Polska jako kraj wysoko rozwinięty cyfryzacyjnie, również podąża za globalnymi trendami. Technologia ta na dobre zagościła w naszej codzienności. Jak wskazują najnowsze analizy raportu badawczego KPMG – „Perspektywa urzędników i instytucji”, otwartość polskiego społeczeństwa na innowacje jest imponująca – już 69% Polaków sięga po narzędzia oparte na sztucznej inteligencji w sposób regularny. Oznacza to, że niemal siedmiu na dziesięciu z nas wdrożyło AI do swoich codziennych zadań, co plasuje Polskę wyraźnie powyżej średniej światowej wynoszącej 66%. Tak głęboka i powszechna adopcja technologii w społeczeństwie w sposób naturalny i nieunikniony przenika również w struktury państwowe, trwale zmieniając oblicze polskiej administracji publicznej. Potwierdzają to szczegółowe badania raportu NASK „AI w e-administracji publicznej – perspektywa urzędników i instytucji”, które dowodzą, że technologiczna transformacja w urzędach nie jest jedynie odległym planem – ona dzieje się na naszych oczach. Dla przeważającej części aparatu państwowego AI to już stały element środowiska pracy:

- Powszechność kontaktu: Aż 77% urzędników ma styczność ze sztuczną inteligencją przynajmniej raz dziennie, a blisko 15% deklaruje, że obcuje z nią niemal nieustannie.
- Praktyczne wykorzystanie GenAI: Niemal połowa badanych urzędników (46%) w ciągu ostatnich sześciu miesięcy aktywnie wykorzystywała narzędzia generatywnej sztucznej inteligencji do realizacji swoich bezpośrednich obowiązków zawodowych. Przekrój stanowisk i wieku: Choć naturalnymi liderami cyfryzacji są najmłodszy pracownicy urzędów (do 34. roku życia), z których 54% aktywnie korzysta z AI, to technologia ta nie jest wyłącznie domeną młodego pokolenia. Rozwiązania te wspierają codzienną pracę aż 60% pracowników na stanowiskach technicznych oraz – co kluczowe dla zarządzania instytucjami – 53% kadry kierowniczej.⁰¹

Analizując powyższe dane z raportów, można wyciągnąć jeden, wyraźny wniosek: pra-

Częstotliwość kontaktu z technologią AI – rodzaj stanowiska i wiek



Wykres 4.2 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, strona 16

cownicy polskiej administracji publicznej są już w pełni gotowi na systemowe wdrożenie rozwiązań opartych na sztucznej inteligencji. Co więcej, to nie jest innowacja, którą trzeba narzucać im odgórnie – oni sami już po nią sięgnęli, dostrzegając jej gigantyczną, realną wartość w codziennych obowiązkach.

Wobec tych twardych danych i statystyk, kształtuje się fundamentalna teza – „Wdrażanie sztucznej inteligencji do sektora publicznego, to nie jest już pytanie „czy”, lecz wyłącznie „kiedy” i „jak szybko to nastąpi”

Powoli dostrzegamy starania ku rozwojowi sztucznej inteligencji w sektorze publicznym. Pierwszym podstawowym fundamentem tej zmiany cyfryzacyjnej jest udostępniony przez rodzime konsorcjum instytucji naukowych duży model językowy PLLuM, który został wdrożony do aplikacji mObywatel. Powstanie

tego modelu potwierdziło zapotrzebowanie na bezpieczną sztuczną inteligencję. Z raportu „Przewodnik po Sztucznej Inteligencji dla Administracji Publicznej” opublikowanym przez Ministerstwo Cyfryzacji, wynika, że pierwsze urzędy już wyraziły chęć realnych wdrożeń. Pierwsze pionierskie wdrożenia w miastach już się odbywają – w Częstochowie wdrożono model PLLuM, który koncentruje się na wsparciu urzędników w ich codziennej pracy, poprzez: streszczanie dokumentów, wyszukiwanie informacji czy przygotowanie pism. W Gdyni natomiast uruchomiono integrację modelu PLLuM z wyszukiwarką Biuletynu Informacji Publicznej (BIP), co umożliwia mieszkańcom wyszukiwanie odpowiednich pism i dokumentów, bezpośrednio z asystentem AI.⁰²

Korzystanie z GenAI w celach zawodowych w ciągu ostatnich 6 miesięcy

Typ urzędu, stanowisko i wiek



Typ urzędu



Rodzaj stanowiska



Wiek



Wykres 4.3 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Obsługa obywateli

Wszystkie te kroki pokazują, że istnieją obszary do zagospodarowania zarówno dla obywateli jak i urzędników. Jednym z najbardziej obiecujących – i najbardziej potrzebnych pól do wdrożeń sztucznej inteligencji jest bezpośrednia obsługa mieszkańców. To tutaj technologia może najszybciej rozładować urzędowe kolejki i zredukować czas obustronnej komunikacji. W kontekście obsługi obywatela wyróżniono dwa kluczowe obszary, w których sztuczna inteligencja jest w stanie realnie wpłynąć na procedury:

1. Bariera językowa i sprawna obsługa

cudzoziemców Polska administracja staje dziś przed dużym wyzwaniem w obliczu sytuacji geopolitycznej. Według Głównego Urzędu Statystycznego (GUS), w ostatnim dniu lipca 2025 r. pracę w Polsce wykonywało 1106,3 tys. cudzoziemców, co oznacza wzrost o 6,2% w stosunku do roku poprzedniego. Obsługa ponad miliona pracowników z zagranicy staje się dużym wyzwaniem dla całego aparatu administracyjnego. Nie bez przyczyny w załączonym poniżej wykresie z badań NASK, wynika, że 77% badanych wskazuje, że to właśnie „Tłumaczenie oraz redakcja treści w językach obcych” zajmuje pierwsze miejsce wśród obszarów do zagospodarowania przez AI. Badania te mają także potwierdzenie w analizach firmy Thunai, zajmującej się rozwiązaniami tłumaczeń z wykorzystaniem AI, według ich danych automatyczne tłumaczenie rozmów może znacząco skracać czas obsługi interesantów w porównaniu do tradycyjnych metod.⁰¹

stron urzędowych i formularzy, to bardzo duże wyzwanie. Zamiast zmuszać obywatela do znajomości i wyszukiwania odpowiednich wniosków, urzędy mogłyby wdrażać rozwiązania, które pomogłyby zadawać pytania w języku naturalnym np. „Jaki wniosek muszę złożyć, aby otrzymać świadczenie 800+?” . Dobrym przykładem jest tutaj wdrożenie wspomniane w poprzednim nagłówku – miasto Gdynia, które zdecydowało się na połączenie możliwości PLLuM z zasobami Biuletynu Informacji Publicznej (BIP), wspomagając i upraszczając przy tym dostęp do informacji publicznych dla swoich mieszkańców.

2. Asystenci obywatela i inteligentne prze-

szukiwanie dokumentów Drugą najwyżej ocenianą pozycją w załączonym wykresie NASK, jest „Wyszukiwanie i analiza istniejącej dokumentacji”. Z perspektywy obywatela, poruszanie się przez gęszcz

W jakim stopniu AI może być wykorzystana w poszczególnych obszarach pracy w urzędzie?



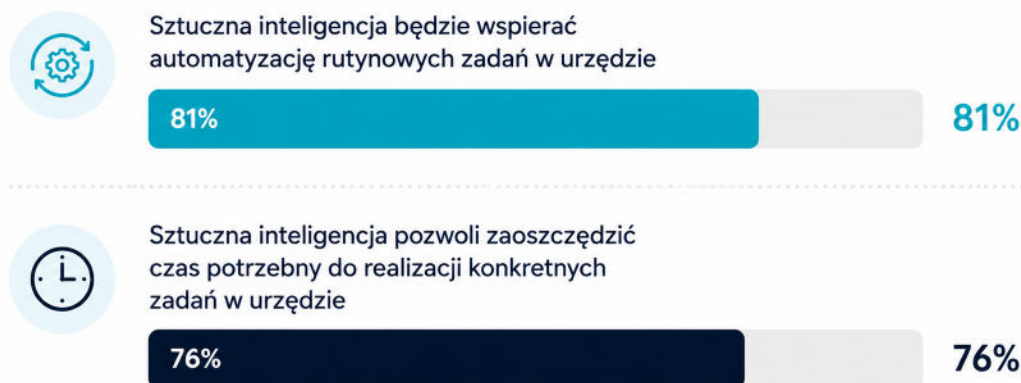
Wykres 4.4 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Choć te innowacyjne podejścia do obsługi klienta stanowią istotną stronę tej technologicznej transformacji cyfrowej, pełny potencjał sztucznej inteligencji ujawnia się dopiero na administracyjnym zapleczu. Ociążenie pierwszej linii kontaktu z obywatelem naturalnie otwiera bowiem drogę do głębokiej optymalizacji procesów wewnętrznych. Płynnie prowadzi to nas do drugiego, fundamentalnego filaru tej rewolucji: bezpośredniego wsparcia codziennej pracy samych urzędników oraz modernizacji urzędowych struktur IT.

Wsparcie urzędników oraz działów IT

Naturalnym i oczywistym faktem jest to, że sztuczna inteligencja nie zastąpi człowieka. Jak w przypadku każdego narzędzia, które zrewolucjonizowało rynek pracy, celem AI jest wsparcie człowieka. Ma na celu przyspieszenie wykonywania obowiązków lub łatwiejsze zrozumienie procesów. Powołując się na badania NASK, wynika, że 81% urzędników jest przekonanych, że AI będzie wspierać automatyzację rutynowych zadań, 76% uważa, że sztuczna inteligencja pozwala efektywnie zaoszczędzić czas.⁰¹

Przewidywany wpływ AI na pracę urzędu



Wykres 4.5 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Pracownicy instytucji publicznych doskonale zdają sobie sprawę z potencjału rozwiązań sztucznej inteligencji. Aż 72% badanych odpowiedziało, że AI ułatwiłoby im pracę z dokumentami w różnych formatach, natomiast aż 67% badanych stwierdziło, że wyszukiwanie informacji przy pomocy języka naturalnego może wspierać ich pracę.⁰²

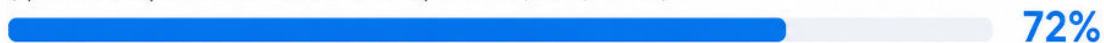
Podobnie dane mają się dla pracowników z obszaru struktur IT. Bowiem aż 82% osób zatrudnionych na stanowiskach technicznych w administracji deklaruje kontakt z technologiami sztucznej inteligencji minimum raz dziennie. Spoglądając na konkretne zastosowania AI w urzędach, specjaliści IT najczęściej

wykorzystują ją do błyskawicznego poszukiwania dokumentacji oraz instrukcji potrzebnych do wykonania zadań technicznych, co stanowi stały element pracy dla 55% z nich. Ponadto, technologia ta aktywnie wspiera sam proces powstawania urzędowego oprogramowania – blisko co trzeci pracownik techniczny (28%) deleguje algorytmom zadania związane z pisaniem, edytowaniem i debugowaniem kodu. Warto również podkreślić, że sztuczna inteligencja pełni w tych zespołach funkcję asystenta edukacyjnego; 32% badanych używa jej do indywidualnego uczenia się i wyjaśniania złożonych zagadnień, co może zwiększać ich produktywność.⁰³

W jakim stopniu wymienione funkcjonalności AI mogą wspierać pracę urzędników?

Praca z dokumentami w różnych formatach

(np. możliwość przeszukiwania i analizowania plików PDF, Word, skanów)



Wyszukiwanie informacji w dokumentach urzędowych za pomocą pytań zadawanych w języku potocznym



Chatboty lub wyszukiwarki dla obywateli pomagające wypełnić wniosek lub znaleźć potrzebne informacje



Automatyczne tworzenie pism urzędowych różnych typów

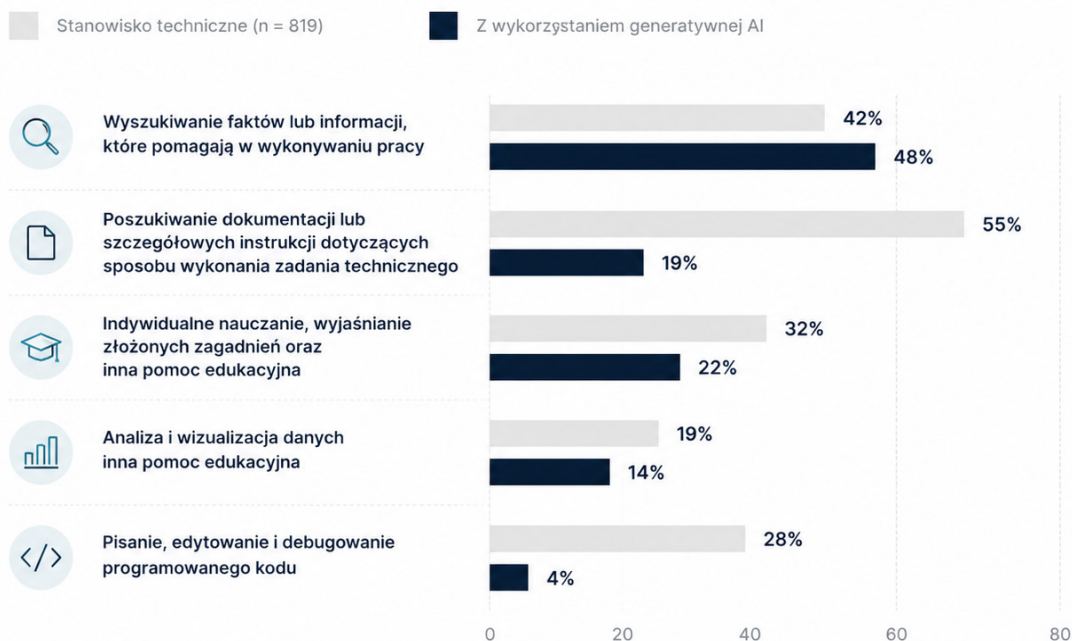
(np. wnioski, decyzje, odpowiedzi na zapytania, notatki służbowe)



Wykres 4.6 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Wykorzystanie GenAI do wykonywania zadań w pracy w ciągu tygodnia przed udziałem w badaniu – stanowiska techniczne i niotechniczne

(5 kategorii charakteryzujących się największymi różnicami w odpowiedziach pomiędzy grupami)



Wykres 4.7 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Istnieje wiele różnych publikacji, które potwierdzają skuteczność użytkowania narzędzi AI w obszarze zagospodarowania IT. Zgodnie z raportem McKinsey „Unleashing developer productivity with generative AI”, sztuczna inteligencja drastycznie zwiększa szybkość wykonywania zadań stricte pro-

gramistycznych, jednocześnie nie obniżając przy tym jakości i bezpieczeństwa tworzenia oprogramowania. Twórcy artykułu dowodzą, że wpływ AI na pracę inżynierów IT jest ogromny, mimo wszystko jednak, ostateczna skala zależy od specyfiki konkretnego zadania.⁰⁴

Generatywna AI może zwiększać wydajność pracy programistów, ale w mniejszym stopniu w przypadku złożonych zadań

Czas realizacji zadań przy użyciu generatywnej AI, %



Wykres 4.8 Źródło: Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025

Jak obrazuje powyższy wykres, największe przyspieszenie widać w obszarach stanowiących żmudną i rutynową część pracy inżyniera IT. Dzięki wsparciu AI, czas potrzebny na realizację poszczególnych etapów drastycznie spada:

- Dokumentowanie kodu: odbywa się o 45–50% szybciej.
- Generowanie nowego kodu: pozwala zaoszczędzić od 35% do 45% czasu, m.in. skutecznie eliminując problem „pustej kartki”.
- Refaktoryzacja kodu: czyli optymalizacja już istniejącego oprogramowania, zajmuje o 20–30% mniej czasu w stosunku do tradycyjnych metod.

Warto także podkreślić, że w procesie wdrażania tych technologii nie można wpadać w bezkrytyczny optymizm. Z danych jednoznacznie wynika, że mit, jakoby sztuczna inteligencja zwalniała programistę z abstrakcyjnego myślenia lub mogła go całkowicie zastąpić, jest daleki od prawdy. Autorzy raportu wskazują bowiem, że przy zadaniach o wysokiej złożoności – na przykład wymagających pracy w nieznanym wcześniej frameworku – realne oszczędności czasu są znikome i spadają poniżej 10%.

Badania wskazują jednoznacznie, że **sztuczna inteligencja w programowaniu, jest tak skuteczna, jak doświadczony i dobry jest korzystający z niej specjalista**. Jest wiele haseł mówiących o zastąpieniu człowieka przez AI, co na chwilę obecną jest dalekie od prawdy. Narzędzia te można w sposób bardzo kolokwialny przyrównać do ekskluzywnego sportowego samochodu – bowiem nawet najlepszy samochód, bez doświadczonego kierowcy rajdowego, staje się nie tylko zagrożeniem, lecz narzędziem, które nie spełnia swojej roli.



Wnioski operacyjne

Sztuczna inteligencja to zarówno przeszłość, teraźniejszość jak i przyszłość. Jako ludzkość zawsze dążymy w kierunku rozwoju i ułatwiania sobie życia. Rozwiązania sztucznej inteligencji już dzisiaj zmieniają wiele w sferze zarówno sektora prywatnego jak i publicznego. Synteza przytoczonych danych badawczych, trendów rynkowych i analiz eksperckich pozwala na wyciągnięcie fundamentalnego wniosku: cyfrowa transformacja polskiej administracji z wykorzystaniem sztucznej inteligencji to już nie jest wizja przyszłości, lecz dynamicznie tocząca się rzeczywistość. W obliczu rewolucji, która w zaledwie kilka miesięcy zdobyła setki milionów użytkowników na całym świecie, aparat państwowy znalazł się w punkcie bez powrotu. Trzeba powiedzieć jasno, jako państwo i obywatele, jesteśmy gotowi, by wdrażać skuteczną, pomocną i bezpieczną sztuczną inteligencję. Aby tego dokonać kolejność wdrażania powinna opierać się obecnie na trzech podstawowych i kluczowych filarach:

1. **Bezwzględny priorytet bezpieczeństwa i zgodności prawnej** - Wdrożenia systemów AI nie mogą odbywać się kosztem bezpieczeństwa. Administracja publiczna operuje na ogromnych zbiorach danych obywateli. Priorytetem powinno być dostosowanie się do regulacji prawnych takich jak RODO i AI Act, by chronić suwerenność i bezpieczeństwo Polek i Polaków.
2. **Konieczność uregulowania oddolnej rewolucji** – Zarówno obywatele, jak i sami urzędnicy są już gotowi na tę zmianę. Brak systemowych, państwowych narzędzi nie zatrzymuje innowacji, lecz spycha ją do szarej strefy. Aby chronić wrażliwe dane obywateli i zachować suwerenność technologiczną, państwo musi natych-

miast dostarczyć urzędnikom bezpieczne, krajowe alternatywy (takie jak model PLLuM czy Bielik), eliminując tym samym ryzyko wycieków informacji do serwerów chmurowych.

3. **Nowa jakość obsługi obywatela** – Administracja musi dogonić standardy wyznaczone przez sektor prywatny. Sztuczna inteligencja staje się kluczem do rozładowania urzędowych kolejek i przełamywania barier komunikacyjnych.

Teza główna pozostaje niewzruszona: dziś nie pytamy już „**czy**” wdrażać AI w administracji publicznej, tylko „**kiedy**” i „**jak**” zrobić to w sposób bezpieczny, systemowy i efektywny.

Rozdział 4

AI w administracji publicznej - przegląd

1. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, s. 14.
2. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, s. 16.

Obsługa obywateli

1. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, s. 34.

Wsparcie urzędników oraz działów IT

1. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, s. 51.
2. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025, s. 44.
3. NASK – PIB, Badanie ilościowe „AI w e-administracji publicznej”, październik–listopad 2025. Baza: osoby korzystające z GenAI w celach zawodowych, n = 2823.
4. McKinsey & Company, Unleashing developer productivity with generative AI, <https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/unleashing-developer-productivity-with-generative-ai>, załącznik nr 1.

05

Bezpieczeństwo danych i AI



krajowa
izba
gospodarki
cyfrowej

Bezpieczeństwo danych i AI – wprowadzenie

Dynamiczny rozwój technologii sztucznej inteligencji w istotny sposób wpływa na funkcjonowanie różnych sektorów gospodarki, w tym administracji publicznej. Rosnąca rola AI w codziennych procesach administracyjnych wymaga od pracowników sektora publicznego dostosowania się do nowych realiów oraz zdobycia odpowiednich kompetencji w tym obszarze.

Sztuczna inteligencja staje się kluczową technologią, która przekształca sposób działania administracji publicznej i wymaga kompleksowego wsparcia w zakresie rozwoju kompetencji. Nowoczesna administracja powinna nie tylko rozumieć potencjał AI, lecz także umiejętnie, świadomie i bezpiecznie wdrażać ją w praktyce.

Sztuczna inteligencja przynosi wiele korzyści – zwiększa efektywność, automatyzuje zadania, wspiera podejmowanie decyzji i umożliwia analizowanie dużych zbiorów danych. Dzięki niej firmy mogą szybciej reagować na zmiany rynkowe, lekarze skuteczniej diagnozować choroby, a użytkownicy korzystać z wygodnych narzędzi ułatwiających codzienne życie. Jednak niewłaściwe korzystanie z AI może prowadzić do poważnych problemów, takich jak błędne decyzje, naruszenia prywatności, stronniczość, a nawet oszustwa i cyberzagrożenia. Aby zminimalizować ryzyko, warto znać najczęstsze błędy związane z wykorzystaniem sztucznej inteligencji oraz sposoby ich unikania.

Jednym z najczęstszych problemów jest brak krytycznego zaufania do wyników generowanych przez AI. Systemy sztucznej inteligencji, mimo swojej zaawansowanej natury, nie są nieomyłne. Działają na podstawie danych, na których zostały wytrenowane, a te mogą być niepełne, nieaktualne lub obciążone błędami. Użytkownicy często traktują odpowiedzi AI jako w pełni wiarygodne, co może prowadzić

do podejmowania niewłaściwych decyzji. Dlatego niezwykle ważne jest, aby zawsze weryfikować uzyskane informacje, szczególnie w kontekście decyzji biznesowych, medycznych czy prawnych.

Kolejnym istotnym zagrożeniem jest naruszenie prywatności. Wprowadzanie wrażliwych danych do systemów AI, takich jak dane osobowe, finansowe czy poufne informacje firmowe, może prowadzić do ich nieautoryzowanego wykorzystania lub wycieku. Wiele narzędzi AI przetwarza dane w chmurze obliczeniowej, co dodatkowo zwiększa ryzyko. Dlatego należy zachować ostrożność i unikać udostępniania informacji, które nie powinny być publicznie dostępne. Organizacje powinny także wdrażać odpowiednie polityki bezpieczeństwa i szkolić pracowników w zakresie ochrony danych.

Stronniczość algorytmów to kolejny poważny problem. AI uczy się na podstawie istniejących danych, które mogą odzwierciedlać uprzedzenia społeczne lub błędy systemowe. W efekcie systemy mogą podejmować decyzje dyskryminujące określone grupy ludzi, na przykład w procesach rekrutacyjnych, kredytowych czy sądowych. Aby temu zapobiegać, konieczne jest stosowanie różnorodnych i reprezentatywnych zbiorów danych oraz regularne audyty systemów AI pod kątem sprawiedliwości i przejrzystości.

Niebezpieczeństwem jest także wykorzystywanie AI do tworzenia fałszywych informacji. Technologie generatywne umożliwiają tworzenie realistycznych tekstów, obrazów czy nagrań wideo, które mogą być używane do manipulacji opinią publiczną lub rozpowszechniania dezinformacji. Tak zwane „deepfakes” mogą podszywać się pod prawdziwe osoby, co stanowi zagrożenie zarówno dla jednostek, jak i dla społeczeństwa jako całości. W związku z tym ważne jest rozwijanie umiejętności krytycznego myślenia oraz korzystanie z wiarygodnych źródeł informacji.

Kolejnym błędem jest brak zrozumienia ograniczeń technologii. AI nie posiada świadomości ani zdolności do rozumienia kontekstu w taki sposób jak człowiek. Może generować odpowiedzi, które brzmią przekonująco, ale są niepoprawne lub pozbawione sensu. Użytkownicy powinni zdawać sobie sprawę z tego, że AI jest narzędziem wspomagającym, a nie zastępującym ludzkie myślenie i odpowiedzialność. Kluczowe decyzje powinny być zawsze podejmowane przez człowieka.

Istotnym problemem jest także brak odpowiednich regulacji i standardów. Szybki rozwój technologii sprawia, że prawo często nie nadąża za zmianami. W rezultacie pojawia-

ją się luki, które mogą być wykorzystywane w nieetyczny sposób. Wprowadzenie jasnych zasad dotyczących projektowania i stosowania AI jest niezbędne, aby zapewnić bezpieczeństwo i ochronę użytkowników. Równie ważna jest odpowiedzialność firm tworzących systemy AI za ich działanie i konsekwencje ich użycia.

Warto również zwrócić uwagę na uzależnienie od technologii. Nadmierne poleganie na AI może prowadzić do osłabienia umiejętności analitycznych i krytycznego myślenia. Użytkownicy mogą przestać samodzielnie rozwiązywać problemy, co w dłuższej perspektywie może mieć negatywne skutki. Dlatego ważne jest zachowanie równowagi między korzystaniem z nowoczesnych narzędzi a rozwijaniem własnych kompetencji.

Sztuczna inteligencja to potężne narzędzie, które może znacząco ułatwić życie i pracę, ale jednocześnie niesie ze sobą liczne zagrożenia. Świadome i odpowiedzialne korzystanie z AI pozwala maksymalizować korzyści przy jednoczesnym minimalizowaniu ryzyka. Kluczem jest edukacja, ostrożność oraz krytyczne podejście do technologii, która – choć niezwykle zaawansowana – nadal pozostaje tylko narzędziem w rękach człowieka.

Komentarz eksperta



Przemysław Kuna

Ekspert

Prezes Urzędu Komunikacji Elektronicznej

Dane dotyczące cyberbezpieczeństwa w Polsce pokazują wyraźnie, że skala zagrożeń dynamicznie rośnie. Jak możemy przeczytać w analizach CERT NASK, w ciągu kilku lat liczba incydentów wzrosła niemal dziewięciokrotnie, a według dostępnych analiz Polska należy dziś do najbardziej atakowanych państw w Unii Europejskiej. Ataki coraz częściej wymierzone są nie tylko w sektor prywatny, ale również w instytucje publiczne, w tym uczelnie, urzędy oraz elementy infrastruktury państwowej.

Mając to na uwadze, szczególnego znaczenia nabiera odporność infrastruktury telekomunikacyjnej, która stanowi fundament funkcjonowania usług cyfrowych, administracji oraz gospodarki. Każde naruszenie bezpieczeństwa sieci lub systemów komunikacyjnych może mieć charakter systemowy i prowadzić do zakłóceń, czy wręcz zagrożeń w wielu obszarach jednocześnie. W tym kontekście bezpieczeństwo danych nie jest wyłącznie zagadnieniem technicznym, lecz elementem stabilności państwa i jego zdolności do świadczenia usług publicznych i zapewnienia bezpieczeństwa obywatelom. Bardzo szybki rozwój zastosowań sztucznej inteligencji dodatkowo zwiększa złożoność tego środowiska. Technologie te umożliwiają automatyzację analizy zagrożeń i mogą wzmacniać systemy bezpieczeństwa, ale jednocześnie tworzą nowe wektory ryzyka. Dotyczy to zarówno skali przetwarzania danych, jak i możliwości wykorzystania AI do prowadzenia działań dezinformacyjnych czy ataków opartych na manipulacji informacją.

W tym kontekście kluczowe znaczenie ma zapewnienie kontroli nad danymi oraz środowiskiem ich przetwarzania. Oznacza to konieczność stosowania rozwiązań gwarantujących przejrzystość operacji, możliwość audytu oraz zgodność z regulacjami. Szczególnie istotne jest ograniczanie przetwarzania danych wrażliwych do środowisk spełniających wysokie standardy bezpieczeństwa, w tym infrastruktury pozostającej pod pełną kontrolą instytucji publicznych lub zaufanych dostawców. Istotnym elementem pozostaje również bezpieczeństwo łańcucha dostaw. Systemy telekomunikacyjne oraz rozwiązania cyfrowe są w dużym stopniu zależne od dostawców technologii, dlatego ich weryfikacja oraz ocena ryzyka stają się kluczowe dla utrzymania integralności infrastruktury i ochrony danych.

Czy sztuczna inteligencja powinna na stałe zagościć w instytucjach publicznych? Odpowiedź jest jednoznaczna – tak. Natomiast bezpieczeństwo danych i zapewnienie ciągłości ich przesyłu to bezpieczeństwo Państwa. Tutaj nie możemy sobie pozwolić, na wprowadzanie nowych technologii w sposób niekontrolowany i pozbawiony procesów lub regulacji, które dadzą możliwość instytucjom państwowym walki z cyberzagrożeniami.

Dane wrażliwe w administracji

Zapewnienie bezpieczeństwa danych w kontekście korzystania ze sztucznej inteligencji jest jednym z najważniejszych obszarów ponieważ korzystanie z ogólnodostępnych rozwiązań AI zarówno płatnych jak i bezpłatnych wiąże się z udostępnieniem danych, które są wprowadzane przez użytkownika. Informacje te mogą zostać wykorzystane do dalszego trenowania modeli, a w niektórych przypadkach również upublicznione. Choć większość dużych przedsiębiorstw technologicznych deklaruje, iż w wersjach komercyjnych zapewnione jest odpowiednie bezpieczeństwo danych, należy uwzględnić fakt, że omawiany obszar technologiczny pozostaje wciąż niewystarczająco uregulowany i monitorowany. W konsekwencji rekomenduje się zachowanie szczególnej ostrożności oraz ograniczonego zaufania podczas pracy z rozwiązaniami opartymi na sztucznej inteligencji. Każdorazowo należy przyjąć założenie, że treści zawarte w zapytaniach (promptach) lub dokumentach przekazywanych do analizy mogą potencjalnie stać się publicznie dostępne.

Podmioty administracji publicznej są szczególnie narażone na ryzyko wycieku danych ponieważ jednym z głównych obszarów ich działalności jest przetwarzanie danych wrażliwych takie jak dane osobowe, adresowe, dane o zdrowiu. W związku z powyższym nie zaleca się umieszczania w zapytaniach następujących kategorii danych i informacji:

- danych osobowych, takich jak rzeczywiste imiona i nazwiska, numery PESEL czy adresy,
- wizerunków osób, w szczególności zdjęć dzieci, klientów lub pracowników,
- informacji objętych tajemnicą handlową, w tym zapisów kontraktowych,
- utworów chronionych prawem autorskim,
- produkcyjnego kodu źródłowego przekazywanego „do poprawy”, zwłaszcza zawierającego hasła, klucze dostępu lub identyfikatory,
- fotografii przedstawiających wnętrza domów lub przestrzeni biurowych.

Bezwzględnie należy przestrzegać wytycznych instytucji lub pracodawcy dotyczących korzystania z określonych narzędzi generatywnej sztucznej inteligencji. W wielu przypadkach rozwiązania udostępniane przez organizacje są uprzednio zweryfikowane, a ich wykorzystanie regulowane jest umowami z dostawcami, zapewniającymi odpowiedni poziom ochrony danych.

W sytuacji korzystania z narzędzi bezpłatnych lub nabytych prywatnie, konieczne jest stosowanie wskazanych zasad oraz uprzednie usuwanie z treści zapytań i dokumentów wszelkich informacji wrażliwych, które mogłyby zostać ujawnione lub niewłaściwie wykorzystane.⁰¹

Przetwarzanie dokumentów

Inteligentne przetwarzanie dokumentów to proces automatycznego wydobywania treści z dokumentów przy wykorzystaniu sztucznej inteligencji. W ostatnich latach zastosowanie AI w tym obszarze wzrosło nawet o 60%.⁰¹

Jednym z kluczowych zastosowań tej technologii jest optyczne rozpoznawanie znaków (OCR). Pozwala ono przekształcać tekst drukowany lub odręczny w formę cyfrową, dzięki czemu staje się on przeszukiwalny i możliwy do edycji. Tradycyjne systemy OCR miały jednak liczne ograniczenia – często nie radziły sobie ze złożonym układem dokumentów, pismem odręcznym ani materiałami o niskiej jakości.

Drugim istotnym zastosowaniem jest inteligentne wyszukiwanie informacji dostępnych w dużych zbiorach danych, które pozwala na szybkie i precyzyjne odnajdywanie potrzebnych treści wśród ogromnej ilości rozproszonych danych. Dzięki wykorzystaniu zaawansowanych algorytmów oraz technik analizy danych możliwe jest nie tylko odnalezienie konkretnych informacji, ale także identyfikowanie powiązań, wzorców oraz zależności między nimi. Tego rodzaju rozwiązania znajdują zastosowanie m.in. w systemach rekomendacyjnych, wyszukiwarkach internetowych, analizie dokumentów czy eksploracji danych biznesowych.

Najważniejsze zagrożenia związane z przetwarzaniem dokumentów przez AI obejmują między innymi:

Jednym z najważniejszych zagrożeń jest ryzyko ujawnienia danych wrażliwych zawartych w dokumentach. Systemy AI często przetwarzają dane osobowe, finansowe lub prawne, które mogą zostać przechwycone przez osoby nieuprawnione lub wykorzystane przez dostawców technologii zwłaszcza

w rozwiązaniach korzystających z ogólnodostępnych modeli AI dostępnych w chmurze. Brak odpowiednich zabezpieczeń lub korzystanie z publicznych narzędzi AI może prowadzić do niekontrolowanego rozpowszechniania informacji.⁰²

Systemy przetwarzające dokumenty mogą zostać zaatakowane zmanipulowanymi danymi wejściowymi lub treningowymi poprzez tak zwane „zatrucie danych” polega na wprowadzaniu spreparowanych informacji do zbiorów treningowych, co prowadzi do błędnych lub celowo zmanipulowanych wyników działania modelu. W praktyce może to oznaczać np. błędną analizę dokumentów finansowych lub prawnych, co stanowi poważne ryzyko dla organizacji.⁰³

Jednym z kluczowych zagrożeń związanych z wykorzystaniem sztucznej inteligencji w przetwarzaniu dokumentów jest ryzyko błędnej interpretacji ich treści. Systemy AI, mimo coraz większego zaawansowania, nie posiadają pełnego zrozumienia kontekstu w takim stopniu jak człowiek, co może prowadzić do nieprawidłowych analiz i wniosków.

Problem ten jest szczególnie widoczny w przypadku dokumentów o złożonej strukturze, takich jak umowy prawne, raporty finansowe czy dokumentacja medyczna. Zawierają one często specjalistyczne słownictwo, wieloznaczne sformułowania oraz odniesienia kontekstowe, które wymagają wiedzy eksperckiej. Modele AI mogą błędnie interpretować znaczenie poszczególnych zapisów, pomijać istotne informacje lub nadawać im niewłaściwy priorytet.

Ryzyka użycia publicznych modeli AI

Korzystanie z publicznych modeli sztucznej inteligencji, takich jak systemy generatywne oparte na dużych modelach językowych (LLM), wiąże się z szeregiem istotnych ryzyk, które mają zarówno charakter techniczny, jak i społeczny oraz indywidualny. Jednym z najważniejszych problemów jest ryzyko dezinformacji. Modele te, mimo wysokiego poziomu zaawansowania, nie posiadają rzeczywistego rozumienia świata, a ich odpowiedzi opierają się na statystycznych wzorcach w danych. W rezultacie mogą generować informacje nieprawdziwe lub nieprecyzyjne, tzw. „halucynacje”, które jednak są przedstawiane w sposób przekonujący. Dla użytkownika oznacza to konieczność każdorazowej weryfikacji uzyskanych informacji, szczególnie w kontekście decyzji wymagających wysokiej dokładności, takich jak kwestie prawne, medyczne czy finansowe.⁰¹

Kolejnym istotnym zagrożeniem jest naruszenie prywatności. Publiczne modele AI często przetwarzają dane wprowadzane przez użytkowników, co może prowadzić do niezamierzonego ujawnienia informacji poufnych. Wprowadzenie do systemu danych osobowych, firmowych czy wrażliwych może skutkować ich dalszym wykorzystaniem w procesach trenowania lub analiz, a w skrajnych przypadkach – ich wyciekiem. Z tego względu szczególnie istotne jest zachowanie ostrożności i unikanie przekazywania informacji, które powinny pozostać poufne.

Modele językowe są również podatne na zjawisko stronniczości (bias), wynikające z charakteru danych, na których zostały wytrenowane. Ponieważ dane te odzwierciedlają istniejące w społeczeństwie uprzedzenia i nierówności, modele mogą je powielać lub nawet wzmacniać. W praktyce może to prowadzić do generowania treści dyskryminujących lub nierzetelnych, co stanowi szczególne zagrożenie w kontekstach edukacyjnych, zawodowych czy społecznych.⁰²

Istotnym problemem jest także brak transparentności działania modeli. Systemy te funkcjonują jako tzw. „czarne skrzynki”, co oznacza, że użytkownik nie ma wglądu w proces podejmowania decyzji ani w źródła informacji, na podstawie których generowana jest odpowiedź. Ogranicza to możliwość oceny wiarygodności uzyskanych wyników oraz utrudnia identyfikację potencjalnych błędów.

Dodatkowo modele AI są podatne na różnego rodzaju ataki, takie jak manipulacja zapytaniami (prompt injection) czy zatrucie danych (data poisoning). Ataki te mogą prowadzić do uzyskiwania nieprawidłowych lub niebezpiecznych odpowiedzi, co stanowi zagrożenie szczególnie w środowiskach profesjonalnych i systemach zautomatyzowanych.

Korzystanie z publicznych modeli AI niesie ze sobą znaczące korzyści, ale również liczne zagrożenia. Kluczowe znaczenie ma świadome i odpowiedzialne podejście użytkownika, oparte na krytycznej analizie informacji, ochronie danych oraz rozumieniu ograniczeń technologii.

Kontrola infrastruktury

Algorytmy sztucznej inteligencji rozwijają się w niezwykle szybkim tempie, osiągając coraz wyższy poziom złożoności i autonomii. Ich zastosowanie obejmuje dziś niemal wszystkie obszary technologii – od analizy danych i automatyzacji procesów biznesowych po systemy bezpieczeństwa i podejmowanie decyzji w czasie rzeczywistym. Jednak wraz z rosnącymi możliwościami AI pojawiają się także nowe, często trudne do przewidzenia zagrożenia dla bezpieczeństwa systemów IT.

Choć sztuczna inteligencja może znacząco zwiększyć efektywność, ograniczyć koszty i poprawić jakość usług, jej implementacja niesie ze sobą ryzyko, którego nie można lekceważyć. W wielu przypadkach zagrożenia te wynikają nie tylko z błędów technicznych, ale również z celowych działań osób trzecich lub niewłaściwego zarządzania systemami.

Jednym z kluczowych problemów są ataki na dane treningowe. Modele AI uczą się na podstawie ogromnych zbiorów danych, dlatego ich jakość ma bezpośredni wpływ na poprawność działania algorytmu. Jeśli cyberprzestępcy zmodyfikują te dane (tzw. „data poisoning”), mogą doprowadzić do sytuacji, w której system zacznie podejmować błędne decyzje lub stanie się podatny na manipulację.

Kolejnym zagrożeniem są decyzje oparte na błędnych założeniach. Algorytmy nie posiadają świadomości ani kontekstu – działają wyłącznie na podstawie dostarczonych danych. Jeśli dane są niekompletne, stronicze lub nieaktualne, AI może generować wyniki, które są pozornie poprawne, ale w rzeczywistości niebezpieczne. W systemach krytycznych, takich jak infrastruktura sieciowa czy systemy finansowe, może to prowadzić do poważnych konsekwencji.

Istotnym problemem jest również utrata prywatności. Systemy AI często przetwarzają ogromne ilości danych osobowych, w tym wrażliwych informacji. Bez odpowiednich zabezpieczeń może dojść do ich wycieku, nie-

autoryzowanego dostępu lub wykorzystania w sposób niezgodny z przeznaczeniem. Dodatkowo, nawet anonimowe dane mogą zostać ponownie zidentyfikowane przy użyciu zaawansowanych technik analizy.

Nadużycia w automatyzacji stanowią kolejne wyzwanie. Automatyczne systemy podejmujące decyzje bez nadzoru człowieka mogą działać w sposób nieprzewidywalny, szczególnie w sytuacjach nietypowych lub kryzysowych. Brak kontroli nad takimi procesami może prowadzić do eskalacji błędów lub podejmowania działań sprzecznych z intencją ich twórców.

W odpowiedzi na te zagrożenia konieczne jest wdrażanie skutecznych mechanizmów zabezpieczających. Jednym z podstawowych kroków, jakie organizacje powinny podjąć, jest przeprowadzanie regularnych audytów bezpieczeństwa. Audyt bezpieczeństwa pozwala na systematyczne identyfikowanie potencjalnych luk w modelach i ich infrastrukturze. Dzięki temu możliwe jest nie tylko wykrycie słabych punktów, ale również zrozumienie, w jaki sposób algorytmy mogą zostać zaatakowane lub zmanipulowane. Tego typu działania powinny obejmować zarówno analizę kodu, jak i testy odporności modeli na różne scenariusze zagrożeń.

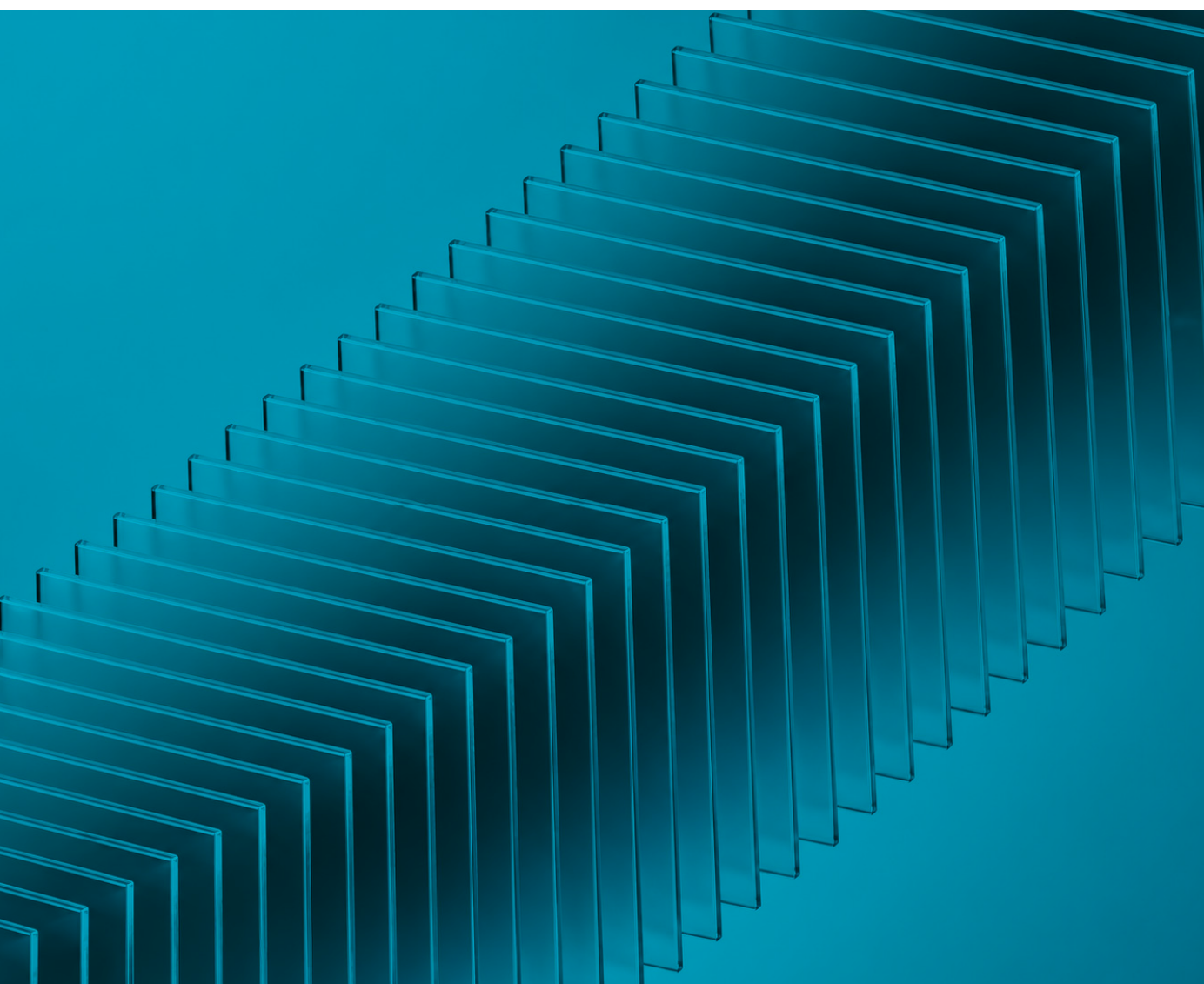
Kolejnym istotnym elementem jest ochrona danych treningowych. Dane stanowią fundament każdego modelu uczenia maszynowego, dlatego ich integralność i jakość mają

bezpośredni wpływ na wyniki działania algorytmu. Wykorzystanie niezabezpieczonych lub niezawerowanych źródeł danych może prowadzić do tzw. „zatrucia danych” (data poisoning), czyli celowego wprowadzania błędnych informacji w celu zaburzenia działania modelu. Organizacje powinny więc stosować mechanizmy kontroli jakości danych, szyfrowanie oraz ograniczenia dostępu.

Nie mniej ważne jest wprowadzenie zasad etyki w obszarze sztucznej inteligencji. Jasno określone wytyczne pomagają nie tylko regulować sposób wykorzystania technologii, ale również zwiększają przejrzystość i zaufanie do systemów AI. Etyczne podejście obejmuje m.in. unikanie uprzedzeń w danych, zapewnienie odpowiedzialności za decyzje podejmowane przez algorytmy oraz ochronę prywatności użytkowników.

Oprócz działań organizacyjnych warto również wdrażać konkretne techniki zabezpieczeń technologicznych. Jednym z najważniejszych zagrożeń są ataki adversarialne, polegające na celowym manipulowaniu danymi wejściowymi w taki sposób, aby model generował błędne wyniki. Obrona przed tego typu atakami wymaga stosowania metod takich jak trening adversarialny, detekcja anomalii czy walidacja danych wejściowych.

Równie istotny jest ciągły monitoring i analiza działania algorytmów. Modele AI funkcjonują w dynamicznych środowiskach, dlatego ich wydajność i bezpieczeństwo mogą się zmieniać w czasie. Regularne monitorowanie pozwala na szybkie wykrycie nieprawidłowości, takich jak nagłe spadki dokładności czy nietypowe wzorce działania, które mogą wskazywać na próbę ataku lub degradację modelu.⁰¹



Komentarz eksperta



Tomasz Panufnik

Dell Technologies

Sales Director Public Sector

Przez lata rozwój IT oznaczał przede wszystkim centralizację – danych, aplikacji i mocy obliczeniowej. Sztuczna inteligencja częściowo odwraca ten kierunek. Im więcej AI wykorzystujemy w codziennych procesach, tym mniej uzasadnione staje się przesyłanie każdego dokumentu czy zapytania do jednego centralnego środowiska. W wielu zastosowaniach efektywniejsze i bezpieczniejsze jest uruchamianie AI możliwie blisko miejsca, w którym znajdują się dane.

Ma to szczególne znaczenie w administracji publicznej. Analiza dokumentu, wyszukanie informacji w aktach czy wsparcie pracownika w przygotowaniu odpowiedzi nie zawsze wymagają dużego modelu działającego w centrum danych.

Nie oznacza to oczywiście, że centrum danych traci znaczenie. Przeciwnie – **powstaje bardziej rozproszona architektura**, w której różne zadania trafiają do różnych warstw infrastruktury. Mały, wyspecjalizowany model może działać na komputerze pracownika, bardziej wymagający proces na lokalnej stacji roboczej lub serwerze, a największe modele i wspólne zasoby danych pozostają w centrum danych. Takie podejście pozwala jednocześnie **ograniczyć przepływ informacji, zmniejszać opóźnienia i lepiej wykorzystywać dostępną moc obliczeniową**.

Jest to istotne również z punktu widzenia **ekonomiki AI**. Koszt uruchamiania modeli językowych szybko spada, a współczesne układy potrafią dostarczać tę samą klasę mocy obliczeniowej przy niższym zużyciu energii niż rozwiązania sprzed kilku lat. Jednocześnie dane, na których pracują organizacje, są coraz bardziej rozproszone – znajdują się w oddziałach, urządzeniach i lokalnych środowiskach, a nie wyłącznie w centralnych centrach danych. Naturalną konsekwencją jest więc przesuwanie części inferencji **AI bliżej tych danych**.

Dlatego kluczowe jest zaprojektowanie architektury, która **pozwała świadomie zdecydować, gdzie dane powinny być przechowywane**, gdzie mogą być przetwarzane i jaka moc obliczeniowa jest rzeczywiście potrzebna dla konkretnego zadania. W administracji publicznej właśnie taka możliwość świadomego rozdzielania obciążeń może stać się jednym z podstawowych elementów bezpiecznej i suwerennej AI.

Audyty i zgodność

Bezpieczeństwo danych w kontekście wdrożenia i funkcjonowania rozwiązań AI stanowi podstawę funkcjonowania każdej nowoczesnej organizacji. Jednym z kluczowych elementów zapewnienia bezpieczeństwa podmiotów przetwarzających dane jest przeprowadzanie audytów bezpieczeństwa IT. Audyt jest szczegółowym procesem polegającym na analizie i ocenie systemów informatycznych funkcjonujących w organizacji. Jego celem jest zapewnienie wysokiej wydajności, odpowiedniego poziomu cyberbezpieczeństwa oraz zgodności zarówno z założeniami biznesowymi, jak i obowiązującymi regulacjami, takimi jak RODO, KSC (Krajowy System Cyberbezpieczeństwa), NIS 2, AI Act, czy ISO/IEC 27001.

Audyty IT umożliwiają organizacjom kompleksową ocenę ich środowiska technologicznego oraz skuteczne zarządzanie ryzykiem i rozwojem infrastruktury. W ramach audytu przeprowadzana jest szczegółowa analiza infrastruktury, systemów i aplikacji, co pozwala na identyfikację słabych punktów, podatności, błędów konfiguracyjnych oraz potencjalnych zagrożeń bezpieczeństwa, a tym samym umożliwia wczesne reagowanie i zapobieganie incydentom. Jednocześnie oceniana jest wydajność oraz efektywność systemów informatycznych, w tym ich skalowalność i wykorzystanie zasobów, co pozwala wykrywać wąskie gardła i lepiej dopasować środowisko IT do potrzeb organizacji. Audyty wspierają także zapewnienie zgodności z obowiązującymi regulacjami prawnymi i standardami branżowymi, minimalizując ryzyko sankcji oraz strat wizerunkowych. Dodatkowo umożliwiają optymalizację kosztów poprzez analizę wydatków na infrastrukturę i usługi IT, w tym outsourcing, co przekłada się na bardziej efektywne zarządzanie zasobami i inwestycjami. Na podstawie wyników audytu możliwe jest również strategiczne planowanie dalszego rozwoju infrastruktury IT z uwzględnieniem przyszłych potrzeb organizacji, postępu technologicznego oraz zmieniającego się krajobrazu zagrożeń.

W odróżnieniu od standardowych metod kontroli w trakcie audytu wykonywana jest szczegółowa i wielopoziomowa analiza

wszystkich komponentów infrastruktury IT. W efekcie uzyskiwany jest pełny obraz poziomu odporności organizacji na zagrożenia cybernetyczne.

Dodatkowo przeprowadzana jest weryfikacja procedur związanych z ochroną danych osobowych oraz informacji poufnych. Dzięki temu możliwe jest spełnienie obowiązujących wymogów prawnych oraz standardów, takich jak RODO, NIS 2. Analizie poddawana jest również podatność pracowników na działania socjotechniczne, co pozwala ocenić czynnik ludzki jako element systemu bezpieczeństwa. W rezultacie ograniczane jest ryzyko wystąpienia poważnych incydentów, w tym wycieków danych.

Istotnym elementem audytu jest raport końcowy, w którym przedstawiane są szczegółowe wnioski oraz rekomendacje działań naprawczych. Wskazywane są konkretne kroki umożliwiające przekształcenie zidentyfikowanych słabości w skuteczne mechanizmy ochronne. W konsekwencji zwiększany jest poziom bezpieczeństwa oraz stabilności funkcjonowania organizacji.⁰¹

Profesjonalnie przygotowany raport z audytu infrastruktury IT powinien być jednocześnie wyczerpujący i czytelny. Jego celem jest jasne przedstawienie wyników analizy oraz zaleceń w taki sposób, aby były zrozumiałe zarówno dla specjalistów IT, jak i osób zarzą-

dzających organizacją. Taki dokument powinien obejmować następujące elementy:

- Podsumowanie dla kierownictwa – najważniejsze wnioski i rekomendacje przedstawione w skróconej formie
- Zakres i metodyka audytu – opis tego, co zostało poddane analizie oraz jakimi metodami się posłużono
- Spis zasobów – dokładna lista sprawdzonych komponentów infrastruktury
- Wnioski – wykryte problemy oraz obszary wymagające poprawy
- Analiza ryzyka – ocena możliwych zagrożeń i ich wpływu na funkcjonowanie organizacji
- Rekomendacje – konkretne sugestie usprawnień uporządkowane według ważności
- Plan wdrożenia – propozycja harmonogramu realizacji zaleceń

Raport z audytu powinien stanowić podstawę do budowy strategii rozwoju i usprawniania infrastruktury IT. Zalecenia należy omówić zarówno w zespole technicznym, jak i na poziomie zarządczym, a następnie przełożyć na konkretny plan działań z jasno określonymi odpowiedzialnościami oraz terminami. Dokument ten warto także traktować jako punkt odniesienia przy kolejnych audytach, co umożliwi ocenę postępów w udoskonalaniu środowiska IT.

Warto pamiętać, że audyt infrastruktury IT to nie jednorazowe zadanie, lecz ciągły proces doskonalenia. Systematyczne jego przeprowadzanie pozwala szybko wykrywać nowe zagrożenia i obszary wymagające optymalizacji, dzięki czemu infrastruktura pozostaje w dobrej kondycji i skutecznie wspiera realizację celów biznesowych.⁰²



Scenariusze ryzyk

Wraz ze wzrostem rozwoju rozwiązań wykorzystujących sztuczną inteligencję, dynamicznie rozwijane są również metody ataku na tego typu systemy. Wynika to z ich rosnącego znaczenia w obszarach takich jak bezpieczeństwo, finanse, medycyna czy infrastruktura krytyczna. Atakujący starają się wykorzystywać specyficzne cechy modeli AI, takie jak podatność na manipulację danymi wejściowymi czy zależność od jakości danych treningowych. Do najpopularniejszych tego typu ataków należą ataki adversarialne, które opierają się na wykrywaniu i wykorzystywaniu słabości modeli AI, aby doprowadzić do błędnych decyzji lub nieprzewidywalnych reakcji. Mogą służyć do destabilizacji pracy systemu, pozyskiwania wrażliwych danych albo uzyskiwania dostępu bez odpowiednich uprawnień. Wraz z rosnącym znaczeniem sztucznej inteligencji i jej zastosowaniem w kluczowych dziedzinach, takich jak finanse, opieka zdrowotna czy bezpieczeństwo publiczne, coraz ważniejsze staje się rozumienie tych ataków oraz rozwijanie skutecznych sposobów ich wykrywania i zapobiegania.

W obszarze cyberbezpieczeństwa wyróżnia się kilka typów ataków adversarialnych, które różnią się zarówno celami, jak i sposobem działania. Każdy z nich stanowi inne zagrożenie dla systemów sztucznej inteligencji. Pierwszy rodzaj to ataki omijające, które koncentrują się na zaburzeniu pracy modelu już po jego wdrożeniu. Ich celem jest wprowadzenie systemu w błąd bez ingerencji w sam model czy dane treningowe. Przykładem może być dodanie niewielkich zmian do obrazu, przez które system rozpoznawania identyfikuje go jako coś zupełnie innego.

Drugą grupę stanowią ataki polegające na wstrzykiwaniu poleceń poprzez manipulowanie danymi wejściowymi tak, aby system wykonał działania, do których nie powinien mieć dostępu. Jest to szczególnie groźne w systemach sterujących procesami przemysłowymi, gdzie może doprowadzić np. do przerwania produkcji.

Kolejny rodzaj ataku stanowi zatrucie modeli AI polegający na celowym zmienianiu danych treningowych, aby zaburzyć działanie modelu. Atakujący mogą wprowadzać spreparowane dane, które powodują, że model generuje błędne lub stronnicze wyniki.

Następną grupą są ataki polegające na ujawnianiu danych, które umożliwiają odtworzenie informacji o danych użytych do trenowania modelu. W przypadku systemów stosowanych np. w medycynie może to prowadzić do odzyskania wrażliwych danych pacjentów i poważnych naruszeń prywatności.

Rozwój metod ochrony przed atakami adversarialnymi stanowi jeden z najważniejszych obszarów współczesnych badań nad bezpieczeństwem systemów sztucznej inteligencji. Istnieje wiele podejść, które mają na celu ograniczenie negatywnych skutków takich ataków, a każde z nich charakteryzuje się własnymi zaletami oraz pewnymi ograniczeniami:

Pierwszą metodą stanowi trenowanie z przykładami zakłóconymi (Adversarial Training), która polega na rozszerzaniu zbioru danych uczących o próbki zawierające celowo wprowadzone zaburzenia (perturbacje). Dzięki temu model uczy się rozpoznawać i uodparniać na tego typu modyfikacje. Wadą tego podejścia jest jednak duży koszt obliczeniowy oraz wydłużony czas procesu treningu.

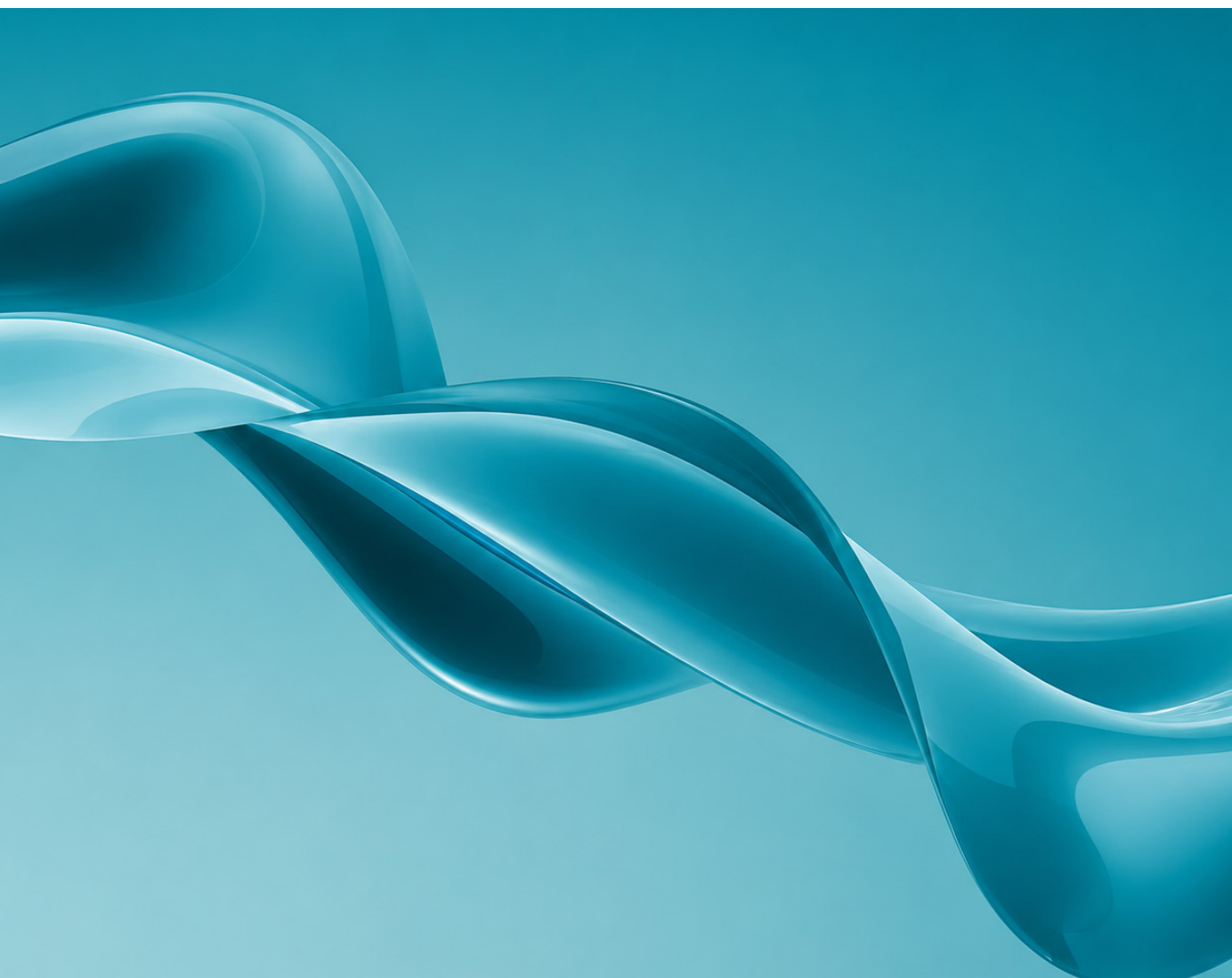
Drugą metodę stanowią techniki wykrywania anomalii, których celem jest identyfikowanie nietypowych lub podejrzanych danych wejściowych, które mogą wskazywać na próbę ataku. W zastosowaniach działających w czasie rzeczywistym, takich jak pojazdy autonomiczne, pozwalają one na szybkie reagowanie na nieprawidłowe lub nieoczekiwane sytuacje.

Kolejną metodę stanowi normalizacja i wstępne przetwarzanie danych polegające na oczyszczaniu danych wejściowych z szumu i nieistotnych elementów, co pomaga modelowi koncentrować się na kluczowych cechach. W rezultacie zmniejsza się jego podatność na subtelne manipulacje adversarialne.

Następną metodą jest uczenie zespołowe (ensemble learning), które opiera się na tre-

nowaniu wielu modeli jednocześnie i łączeniu ich wyników w końcowej decyzji. Takie podejście zwiększa odporność systemu, ponieważ błędy jednego modelu mogą zostać wychwycone lub zrównoważone przez pozostałe.

Ataki adversarialne należą do najpoważniejszych problemów, z jakimi mierzą się dzisiejsze systemy sztucznej inteligencji. Kluczowe znaczenie ma poznanie sposobów ich działania, opracowywanie efektywnych mechanizmów obronnych oraz wdrażanie nowych rozwiązań w zakresie cyberbezpieczeństwa. Wraz z rosnącą obecnością AI w różnych dziedzinach życia, coraz ważniejsze staje się zagwarantowanie jej bezpieczeństwa i stabilności. Aby sprostać tym wyzwaniom, konieczna jest współpraca naukowców, inżynierów i twórców, która umożliwi rozwój sztucznej inteligencji w sposób odpowiedzialny, bezpieczny i zrównoważony.⁰¹



Wnioski bezpieczeństwa

Dynamiczny rozwój sztucznej inteligencji (AI) fundamentalnie zmienia sposób funkcjonowania nowoczesnej gospodarki oraz administracji publicznej. Pomimo tego, że technologie te oferują bezprecedensowy wzrost efektywności i automatyzację złożonych procesów, ich wdrażanie niesie ze sobą istotne wyzwania w zakresie bezpieczeństwa i ochrony danych. Aby w pełni wykorzystać potencjał AI, niezbędne jest świadome i odpowiedzialne podejście do tej technologii.

Jednym z najpoważniejszych błędów użytkowników jest nadmierna wiara w nieomyślność systemów AI. Należy pamiętać, że modele te operują na danych historycznych, które mogą być niepełne, nieaktualne lub zawierać błędy. Zjawisko tzw. „halucynacji”, czyli generowania nieprawdziwych informacji w sposób bardzo przekonujący, sprawia, że każda odpowiedź algorytmu powinna podlegać ludzkiej weryfikacji, szczególnie w obszarach o wysokim stopniu ryzyka, takich jak medycyna, prawo czy finanse.

Administracja publiczna, ze względu na charakter przetwarzanych danych, jest szczególnie narażona na ryzyko ich wycieku. Korzystanie z publicznych, ogólnodostępnych narzędzi AI wiąże się z ryzykiem upublicznienia wprowadzanych treści, które mogą zostać wykorzystane przez właściciela rozwiązania w różnych celach. W związku z tym, użytkownicy powinni bezwzględnie unikać wprowadzania do zapytań (promptów) danych wrażliwych takich jak dane osobowe obejmujące imiona, nazwiska, numery PESEL czy dane adresowe lub wizerunki osób, w tym zdjęcia pracowników oraz klientów a także dane dotyczące tajemnic handlowych, zapisów kontraktowych oraz utworów chronionych prawem autorskim.

Sztuczna inteligencja wprowadza specyficzne rodzaje ataków, na które tradycyjne systemy IT nie zawsze są przygotowane. Do najistotniejszych należą ataki adwersarialne, polegające na celowej manipulacji danymi wejściowymi, zatrutowaniu danych poprzez ingerencje w zbiór treningowy na etapie uczenia modelu, mającą na celu trwałe uszkodzenie jego logiki działania, wstrzykiwania promptów polegającej na manipulowaniu zapytaniami w celu zmuszenia AI do wykonania niedozwolonych działań lub ujawnienia poufnych informacji.

Skuteczna ochrona przed tymi zagrożeniami wymaga systematycznego podejścia. Kluczowym elementem jest przeprowadzanie regularnych audytów bezpieczeństwa, które pozwalają na identyfikację słabych punktów w infrastrukturze oraz pozwolą na zapewnienie zgodności z regulacjami takimi jak RODO, AI Act czy NIS 2. Ważne przy tym jest stosowanie się do zaleceń bezpieczeństwa stanowiących wynik audytu oraz przeprowadzanie audytów w sposób cykliczny.

Rozdział 5

Dane wrażliwe w administracji

1. Gov.pl, Pułapki związane z wykorzystywaniem sztucznej inteligencji – jak unikać zagrożeń, <https://www.gov.pl/web/baza-wiedzy/pulapki-zwiazane-z-wykorzystywaniem-sztucznej-inteligencji--jak-unikac-zagrozen>, dostęp: 25.03.2026.

Przetwarzanie dokumentów

1. McKinsey & Company, The State of AI in 2022 and a Half Decade in Review, <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-in-2022-and-a-half-decade-in-review>, dostęp: 25.03.2026.

2. Zintegrowana Platforma Edukacyjna (ZPE), Zagrożenia przy wdrażaniu rozwiązań AI w edukacji, <https://zpe.gov.pl/a/42-zagrozenia-przy-wdrazaniu-rozwiazan-ai-w-edukacji/DL5knFNZ2>, dostęp: 25.03.2026.

3. Gov.pl, Zagrożenia dla systemów wykorzystujących AI, <https://www.gov.pl/web/baza-wiedzy/zagrozenia-dla-systemow-wykorzystujacych-ai>, dostęp: 25.03.2026.

Ryzyka użycia publicznych modeli AI

1. Gov.pl, Pułapki związane z wykorzystywaniem sztucznej inteligencji – jak unikać zagrożeń, <https://www.gov.pl/web/baza-wiedzy/pulapki-zwiazane-z-wykorzystywaniem-sztucznej-inteligencji--jak-unikac-zagrozen2>, dostęp: 25.03.2026.

Kontrola infrastruktury

1. Owocne Inwestycje, Wpływ AI na bezpieczeństwo systemów IT, <https://owocneinwestycje.pl/wplyw-ai-na-bezpieczenstwo-systemow-it>, dostęp: 25.03.2026.

Audyt i zgodność

1. Neuros, Audyt infrastruktury IT, <https://neuros.pl/audyt-infrastruktury-it/>, dostęp: 25.03.2026.

2. 4hfix, Co powinien zawierać audyt infrastruktury IT, <https://4hfix.pl/co-powinien-zawierac-audyt-infrastruktury-it>, dostęp: 25.03.2026.

Scenariusze ryzyk

1. drmalinowski.edu.pl, Adwersarialne ataki na sztuczną inteligencję, <https://www.drmalinowski.edu.pl/post-s/2824-adwersarialne-ataki-na-sztuczna-inteligencje>, dostęp: 25.03.2026.

06

Studium Przypadków



krajowa
izba
gospodarki
cyfrowej

Case Studies - Safe VOX

Polska administracja publiczna, służba zdrowia i instytucje porządku publicznego stają dziś przed rosnącym wyzwaniem komunikacyjnym. Napływ cudzoziemców - pracowników, uchodźców, studentów i turystów - sprawia, że coraz więcej codziennych interakcji odbywa się ponad barierą językową. Urzędnicy, lekarze i funkcjonariusze muszą sprawnie obsługiwać interesantów, którzy nie posługują się językiem polskim, często w sytuacjach wymagających precyzji, poufności i pełnej dokumentacji przebiegu rozmowy. Tradycyjne rozwiązania takie jak tłumacze przysięgli czy aplikacje konsumenckie nie spełniają wymogów instytucjonalnych, są kosztowne i trudno dostępne w trybie natychmiastowym.

Odpowiedzią na te potrzeby są specjalizowane narzędzia AI do tłumaczenia, zaprojektowane z myślą o środowiskach instytucjonalnych. Działające lokalnie, bez przysyłania danych do zewnętrznej chmury, umożliwiając **transkrypcję, dwukierunkowe tłumaczenie mowy i tekstu w czasie rzeczywistym, automatyczny zapis rozmów oraz ich eksport do akt sprawy, transkrypcję nagrań - plików audio**. Poniżej przedstawiamy przykłady konkretnych zastosowań tego typu rozwiązań w instytucjach: urząd administracyjny, ubezpieczenia społeczne oraz placówki medyczne. Przykładowe zastosowania narzędzia:

Urząd administracyjny - legalizacja pobytu

Cudzoziemiec prowadzi rozmowę dotyczącą legalizacji pobytu, która może mieć bezpośredni wpływ na decyzję administracyjną. Rozwiązanie AI do tłumaczenia precyzyjnie tłumaczy pytania i wyjaśnienia dotyczące statusu pobytu. Wrażliwe dane osobowe wprowadzane są dyskretnie - na ekranie widocznym wyłącznie dla interesanta. Zapis rozmowy w archiwum systemu dokumentuje przebieg całego postępowania, co jest istotne zarówno dla urzędnika, jak i dla samego cudzoziemca. Rozwiązanie AI do tłumaczenia wspiera urzędnika, tłumacząc rozmowę zarówno mowie, jak i piśmie. Dzięki automatycznemu tłumaczeniu i zapisowi rozmowy komunikacja przebiega sprawniej, a czas obsługi sprawy ulega skróceniu.

Ubezpieczenia społeczne - obsługa wniosku o zasiłek chorobowy

Cudzoziemiec składa wniosek o zasiłek chorobowy, nie znając języka polskiego ani obowiązujących procedur. Rozwiązanie AI do tłumaczenia wspiera pracownika instytucji ubezpieczeń społecznych poprzez tłumaczenie rozmowy w czasie rzeczywistym - wyjaśniane są obowiązujące procedury i wytyczne w języku interesanta. System automatycznie archiwizuje przebieg rozmowy, co stanowi potwierdzenie przekazanych informacji i może służyć jako dokumentacja w razie kontroli. Treść rozmowy można wyeksportować bezpośrednio do akt sprawy.

Placówki medyczne - nagłe zdarzenie medyczne

Cudzoziemiec trafia na szpitalny oddział ratunkowy z nagłym problemem zdrowotnym. Rozwiązanie AI do tłumaczenia wspiera personel medyczny w prowadzeniu wywiadu w języku pacjenta, co jest kluczowe w sytuacji zagrożenia życia. System zapisuje informacje przekazane przez pacjenta - w tym zgody na leczenie i ustalenia medyczne. Wrażliwe dane pacjent może wprowadzać samodzielnie i dyskretnie na własnym ekranie, co chroni jego prywatność nawet w przestrzeni publicznej oddziału.

Komentarz eksperta



Magdalena Gaj

Ekspert

Członkini Rady Nadzorczej Exatel, była Prezes Zarządu Exatel

Raport bardzo dobrze pokazuje, że suwerenność technologiczną należy dziś analizować nie tylko w kontekście rozwoju IT, ale przede wszystkim przez pryzmat bezpieczeństwa państwa i klasycznej triady CIA – poufności, integralności oraz dostępności danych i usług. Właśnie te trzy elementy stają się fundamentem nowoczesnej odporności cyfrowej państwa.

Istotne jest to, że raport nie utożsamia suwerenności technologicznej z izolacją od globalnych technologii. Nie chodzi o budowanie wszystkiego samodzielnie, tylko o zachowanie realnej kontroli nad danymi, infrastrukturą i ciągłością działania państwa w sytuacjach kryzysowych.

Bardzo ważna jest teza dotycząca suwerennej chmury. Dla administracji publicznej nie jest to dziś kwestia wygody, czy modelu zakupowego, ale wymóg bezpieczeństwa. Państwo musi mieć pewność, kto kontroluje dane, gdzie są one przetwarzane i jakie regulacje prawne mają do nich zastosowanie. W tym kontekście raport zwraca uwagę na ryzyko uzależnienia od jednego dostawcy technologii. Vendor lock-in to nie tylko problem biznesowy – to realne ograniczenie odporności i sprawczości państwa.

Raport identyfikuje wyzwania samorządów, które przetwarzają ogromne ilości danych obywateli, często bez wystarczających zasobów i kompetencji do samodzielnej oceny ryzyk technologicznych. Dlatego potrzebne są wspólne standardy, bezpieczna infrastruktura i zaufani partnerzy technologiczni wspierający administrację w procesie transformacji cyfrowej.

Sama suwerenna chmura nie wystarczy, jeśli dane będą przetwarzane przez zewnętrzne modele AI. Chmura i lokalne AI powinny być traktowane jako dwa elementy jednej architektury bezpieczeństwa cyfrowego państwa.

Polska zna już koszt uzależnienia energetycznego, więc nie powinna czekać, aż pozna koszt uzależnienia technologicznego. To jest fundament myślenia o cyfrowym bezpieczeństwie państwa. „Suwerenna chmura”, lokalne modele językowe i krajowa infrastruktura obliczeniowa nie są trzema oddzielnymi tematami. Są to trzy warstwy jednej architektury bezpieczeństwa cyfrowego państwa.

Polska nie wygra z globalnymi gigantami skalą finansowania. Ale możemy wygrać tam, gdzie mamy przewagę: w języku polskim, administracji publicznej, specyfice prawa, danych państwowych, bezpieczeństwie sektorów regulowanych i infrastrukturze krytycznej. Musimy działać mądrze, selektywnie i strategicznie – nie kopiować Doliny Krzemowej, tylko budować własny model odporności cyfrowej.

Suwerenność technologiczna to nie hasło polityczne. To warunek ciągłości państwa, a nawet jego niepodległości. To zdolność do utrzymania poufności danych, integralności procesów i dostępności usług publicznych wtedy, kiedy najbardziej ich potrzebujemy. W czasach wojny hybrydowej, presji informacyjnej, cyberataków i uzależnień infrastrukturalnych nie możemy traktować chmury, AI i łączności jako zwykłych usług IT. To są elementy bezpieczeństwa narodowego.

Case Studies - Safe TAlA

Rosnąca liczba spraw, dokumentów i regulacji prawnych sprawia, że instytucje publiczne oraz podmioty ochrony zdrowia coraz częściej poszukują narzędzi, które pozwolą usprawnić codzienną pracę z dokumentacją. Pracownicy urzędów, szpitali czy instytucji ubezpieczeniowych spędzają znaczną część czasu na ręcznym przeszukiwaniu akt, analizowaniu pism i przygotowywaniu odpowiedzi opartych na złożonych podstawach prawnych. Tradycyjne metody pracy stają się niewystarczające wobec rosnących wymagań proceduralnych i oczekiwań dotyczących szybkości obsługi.

Odpowiedzią na te potrzeby są narzędzia AI do analizy danych, które umożliwiają automatyczne przetwarzanie dużych zbiorów dokumentów, inteligentne wyszukiwanie informacji z zachowaniem kontekstu zapytania oraz wsparcie w przygotowywaniu decyzji i odpowiedzi formalnych. Rozbudowa narzędzia konwersacyjnego pozwala na tworzenie i uruchamianie skryptów, np. do przetwarzania arkuszy kalkulacyjnych, natomiast integracja z zewnętrznymi systemami, takimi jak SSO, SharePoint czy EKD, umożliwia automatyzację procesów. Kluczowym wymogiem w sektorze publicznym i ochronie zdrowia jest bezpieczeństwo danych wrażliwych - przetwarzanie informacji wyłącznie w lokalnej infrastrukturze instytucji, bez przesyłania ich do zewnętrznych systemów chmurowych, staje się standardem zgodnym z wymogami RODO, NIS 2 oraz krajowych regulacji bezpieczeństwa.

Poniżej przedstawiono przykładowe przypadki użycia narzędzi AI do analizy danych w wybranych instytucjach.

Urząd administracyjny - analiza dokumentacji postępowania administracyjnego

Pracownik urzędu analizuje dokumenty w ramach postępowania administracyjnego dotyczącego cudzoziemca. Narzędzie AI do analizy danych umożliwia analizę i porównanie nadesłanych dokumentów z aktami sprawy,

a następnie pomaga przygotować propozycję decyzji administracyjnej. Całość przetwarzana jest lokalnie, bez przesyłania danych na zewnątrz. Przetwarzanie archiwalnych akt spraw: pracownik analizuje archiwalne akta w kontekście konkretnej sprawy. Narzędzie AI do analizy danych wspiera go w szybkim przetwarzaniu dużych wolumenów dokumentów, skutecznym wyszukiwaniu informacji z zachowaniem kontekstu pytania oraz w formułowaniu końcowych analiz i wniosków. System jest szczególnie przydatny przy pracochłonnym przeszukiwaniu wielostronicowych akt.

Ubezpieczenia społeczne - analiza wniosków świadczeniowych

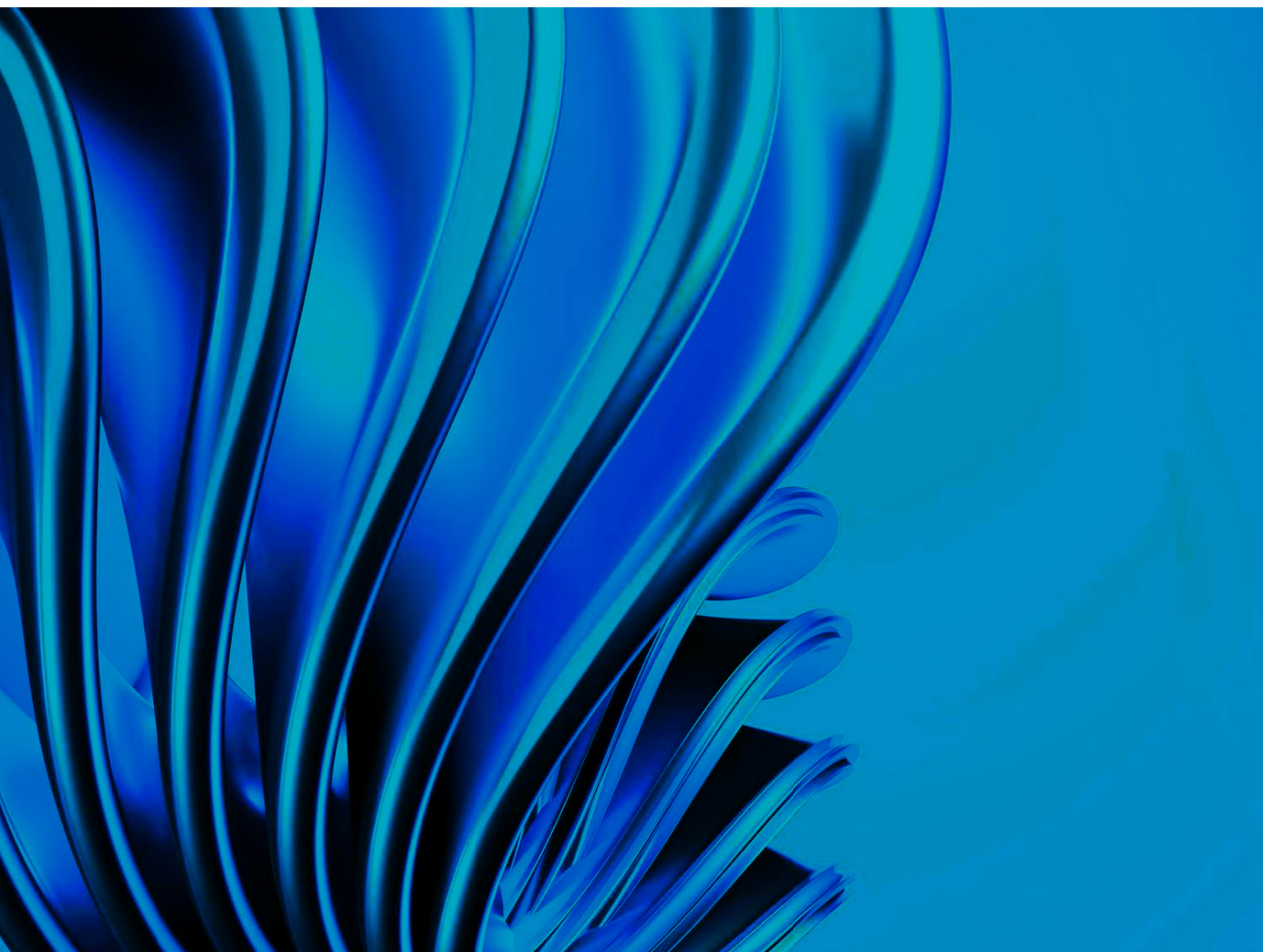
Pracownik odpowiada na wnioski świadczeniowe, które wymagają uwzględnienia obowiązujących przepisów, interpretacji i wewnętrznych wytycznych. Narzędzie AI do analizy danych wspiera go w przygotowaniu formalnej odpowiedzi zgodnej z przepisami, pomaga zachować spójność językową i proceduralną, a także generuje projekt odpowiedzi zgodny z obowiązującą procedurą.

Placówki medyczne - analiza wrażliwej dokumentacji medycznej, reklamacje

Pracownik placówki medycznej analizuje dokumentację medyczną wymagającą

bezpiecznego przetworzenia. Narzędzie AI do analizy danych wyjaśnia treść dokumentacji, ogranicza ryzyko błędnej interpretacji danych medycznych oraz zapewnia dostęp do danych wrażliwych wyłącznie uprawnionym użytkownikom. Ze względu na wrażliwy charakter danych kluczowe jest tu przetwarzanie wyłącznie w lokalnej infrastrukturze, bez kontaktu z chmurą. Narzędzie AI pozwala w inteligentny sposób przeszukiwać zasoby wiedzy szpitala pod kątem konkretnych zapisów proceduralnych i regulaminowych. System umożliwia personelowi zadawanie pytań w języku naturalnym i błyskawiczne odnajdywanie odpowiednich fragmentów dokumentacji administracyjnej, wraz ze wskazaniem źródła i kontekstu danego zapisu. Dodatkowym elementem jest funkcja automatycznej analizy konkretnych przypadków medycznych względem obowiązujących procedur, co wspiera personel w szybkiej i wiarygodnej ocenie zgodności podejmowanych działań z wewnętrznymi standardami placówki.

Zastosowanie: rozpatrywanie skarg pacjentów - szybka weryfikacja zgodności działań personelu z regulaminem i procedurami wewnętrznymi, z odwołaniem do konkretnego źródła zapisu. Rozwiązania istotnie skracają czas analizy dokumentacji administracyjnej i medycznej, eliminując konieczność ręcznego przeszukiwania obszernych zbiorów procedur. Personel szpitala zyskuje nowoczesne narzędzia ułatwiające zarządzanie zawartością archiwum dokumentów, co pozwala ograniczyć czas poświęcany na czynności administracyjne i skierować większą uwagę na zadania kliniczne oraz bezpośredni kontakt z pacjentem. Automatyzacja analizy zgodności proceduralnej przekłada się także na większą spójność i wiarygodność podejmowanych decyzji, ograniczając ryzyko pominięcia istotnych zapisów regulaminowych.



Digitalizacja archiwum dokumentacji w instytucji publicznej

Wyzwanie

Instytucja publiczna gromadzi na przestrzeni lat obszerną dokumentację papierową - akta spraw, wnioski, decyzje administracyjne i inne materiały archiwalne. Wraz z upływem czasu ich liczba osiąga skalę utrudniającą sprawne zarządzanie: dokumentacja zajmuje znaczną powierzchnię magazynową, a odnalezienie konkretnego dokumentu wymaga ręcznego przeszukiwania archiwum przez personel. W praktyce przekłada się to na wydłużony czas obsługi spraw, większe ryzyko błędów oraz obciążenie pracowników zadaniami o niskiej wartości dodanej, zamiast pracy merytorycznej z interesantem.

Rozwiązanie

Wdrożone narzędzie AI automatyzuje proces digitalizacji dokumentacji papierowej. System skanuje i rozpoznaje treść oraz strukturę dokumentów, a personel ogranicza się do weryfikacji i zatwierdzenia wyniku, zanim dokument trafi do centralnego systemu informatycznego instytucji. Zdigitalizowany zasób można następnie przeszukiwać w sposób inteligentny - w języku naturalnym, bez konieczności znajomości struktury archiwum czy fizycznego przeglądania akt. Rozwiązanie działa w infrastrukturze instytucji, co zapewnia pełną kontrolę nad danymi oraz zgodność z wymogami ich ochrony. Wdrożenie i dostosowanie do specyfiki placówki przebiega etapowo, bez konieczności wstrzymywania bieżącej pracy archiwum.

Rezultaty

Czas dostępu do dokumentacji ulega istotnemu skróceniu, a personel zostaje odciążony od czynności manualnych na rzecz zadań wymagających bezpośredniego kontaktu z interesantem. Digitalizacja porządkuje i zabezpiecza zasoby archiwalne, ograniczając ryzyko utraty lub uszkodzenia dokumentów oraz zmniejszając zapotrzebowanie na fizyczną przestrzeń magazynową.

Komentarz eksperta



dr inż. Karol Antczak

CTO Safe Stack

Pracownik naukowy w Wojskowej Akademii Technicznej

Obecnie wchodzimy w erę coraz surowszych regulacji dotyczących sztucznej inteligencji i cyberbezpieczeństwa. Kluczowe akty prawne to: AI Act (Rozporządzenie o Sztucznej Inteligencji), dyrektywa NIS 2 (UKSC) oraz DORA (Digital Operational Resilience).

AI Act wprowadza podejście „oparte na ryzyku”: systemy o wysokim lub niedopuszczalnym ryzyku (np. masowe rozpoznawanie twarzy, tzw. „scoring” społeczny, podejmowanie decyzji wobec obywateli) są z góry wyłączone albo muszą być poddane szczególnej kontroli, a publiczni administratorzy jako „użytkownicy wysokiego ryzyka” muszą m.in. prowadzić rejestr używanych systemów AI, dokumentację techniczną i zapewniać nadzór ludzki oraz prowadzić szkolenia.

NIS 2 i DORA natomiast mocniej wiążą cyberbezpieczeństwo z całością zarządzania: błędy w bezpieczeństwie stały się problemem nie tylko działów IT, ale także zarządu, rady nadzorczej i managementu, na które mogą być nałożone bardzo wysokie kary finansowe za nieprzestrzeganie przepisów i niedokładanie należytej staranności.

Dla urzędów i samorządów nowe przepisy oznaczają przejście z „projektów pilotażowych z AI” do operacyjnego, nadzorowanego stosowania technologii AI. Rozporządzenie wymusza scedowanie odpowiedzialności: bez względu na dostawcę, gmina czy resort musi sprawdzić, czy system nie łamie pewnych zasad, np. nie śledzi obywateli w przestrzeni publicznej ani samoistnie nie podejmuje krytycznych decyzji, które, zgodnie z prawem ma podjąć urzędnik, a nie automat.

NIS 2 z kolei objął urzędy publiczne, które mają obowiązek prowadzenia regularnych analiz ryzyka, testów bezpieczeństwa, szkoleń, wdrażania planów reagowania na incydenty, raportowania poważnych naruszeń do krajowych organów oraz wdrożenia systemu monitorowania (SIEM / SOC L1) i reagowania (np. SOC L2, XDR, SOAR).

Kluczem do zminimalizowania ryzyka jest połączenie silnego zarządzania ryzykiem, odpowiedzialnego podejścia do narzędzi AI oraz nowoczesnych zabezpieczeń cyfrowych, tak aby urząd działał sprawnie i szybko – ale nie kosztem bezpieczeństwa i zaufania obywateli.

Case Studies - Safe CODER

Cyfryzacja administracji publicznej w Polsce i Europie nabrała w ostatnich latach tempa, którego systemy informatyczne instytucji publicznych często nie nadążają obsłużyć. Kolejne regulacje - od nowelizacji ustaw krajowych, przez wdrożenia wymagań unijnych, po obowiązki wynikające z NIS 2 czy AI Act generują nieustanną presję na działy IT, które muszą rozwijać, aktualizować i dokumentować oprogramowanie szybciej niż dotychczas, przy jednoczesnym zachowaniu najwyższych standardów bezpieczeństwa i zgodności z prawem. W tym samym czasie rynek pracy programistów pozostaje jednym z najbardziej napiętych w całej gospodarce. Publiczne działy IT konkurują o specjalistów z sektorem prywatnym, rzadko mogąc zaoferować porównywalne wynagrodzenia. Efektem są chroniczne niedobory kadrowe, rosnące uzależnienie od zewnętrznych podwykonawców oraz trudności z utrzymaniem ciągłości wiedzy o własnych systemach - szczególnie tych starszych, pozbawionych aktualnej dokumentacji technicznej. Odpowiedzią na te wyzwania stają się narzędzia sztucznej inteligencji wspierające wytwarzanie oprogramowania.

Instytucja publicznej opieki zdrowotnej - aktualizacja medycznych systemów IT

Programista zespołu utrzymania systemów medycznych przygotowuje do wdrożenia aktualizację algorytmu rozliczania świadczeń medycznych. Przed przekazaniem zmian na środowisko produkcyjne konieczny jest szczegółowy przegląd bezpieczeństwa kodu. Narzędzie AI wspierające wytwarzanie oprogramowania automatycznie identyfikuje miejsca potencjalnie narażone na podatności oraz generuje dokumentację techniczną gotową do przekazania do działu prawnego.

Administracja publiczna - rozbudowa platformy e-usług

Programista w wydziale informatyki ma za zadanie zintegrować nowy moduł obsługi wniosków budowlanych z istniejącą platformą e-usług. Narzędzie AI wspierające wytwarzanie oprogramowania pozwoli mu w krótkim czasie zrozumieć strukturę kodu oprogramowania, wygenerować szkielet nowego modułu spójny z dotychczasowymi konwencjami projektu oraz przygotować testy jednostkowe weryfikujące poprawność integracji. Dzięki temu urząd nie musi zlecać prac zewnętrznej firmie i dostarcza nową e-usługę terminowo. Analogiczne podejście ma istotne znaczenie w sektorze obronności oraz w instytucjach odpowiedzialnych za bezpieczeństwo państwa i służby mundurowe, gdzie możliwość samodzielnego, szybkiego i w pełni kontrolowanego wytwarzania oraz integracji oprogramowania - bez udziału podmiotów zewnętrznych - ma kluczowe znaczenie z punktu widzenia bezpieczeństwa informacji, niezależności technologicznej i ochrony danych o znaczeniu strategicznym.

Bankowość - systemy transakcyjne, scoring, anti-fraud

Programista pracuje na kodzie krytycznym biznesowo - odpowiadającym za obsługę transakcji, ocenę scoringową klientów oraz mechanizmy anti-fraud w kluczowym systemie bankowym. Ze względu na charakter danych i regulacje sektorowe, kod ten nie może być przesyłany do zewnętrznych modeli AI, co ogranicza możliwość korzystania

z popularnych narzędzi wspierających programowanie. Rozwiązanie Safe CODER pozwala programiście na bezpieczne wykorzystanie AI do analizy i rozwoju kodu bez wysyłania go poza infrastrukturę banku. Narzędzie skraca cykl wytwarzania oprogramowania, automatycznie wskazuje potencjalne podatności bezpieczeństwa oraz wspiera przegląd zmian pod kątem zgodności z regulacjami sektora finansowego.

07

**Suwerenność
technologiczna**

Komentarz eksperta



Adam Paczusi

Pełnomocnik Zarządu ds. Sztucznej Inteligencji w KIGC

CEO i współzałożyciel Safe Stack

Rynek technologii dla sektora publicznego obserwuję od ośmiu lat i przez ten czas widziałem wiele rozwiązań, które świetnie sprawdzały się w prezentacjach, ale zawodziły w realnym wdrożeniu w administracji. Sztuczna inteligencja w instytucjach publicznych rządzi się swoimi prawami – nie wystarczy, że narzędzie jest nowoczesne, musi być przede wszystkim bezpieczne, przewidywalne i zgodne z obowiązującymi przepisami prawa.

W Safe Stack od początku projektowaliśmy nasze produkty z myślą o tych wymaganiach. Nie adaptujemy rozwiązań stworzonych dla biznesu do potrzeb urzędów – zamiast tego tworzymy narzędzia, które od podstaw odpowiadają na specyfikę sektora publicznego: wysokie standardy bezpieczeństwa danych, pełną kontrolę nad informacją oraz zgodność z regulacjami takimi jak RODO, ustawa o krajowym systemie cyberbezpieczeństwa czy wymogi wynikające z AI Act.

Kluczowym wyróżnikiem naszych rozwiązań jest fakt, że działają one lokalnie, w infrastrukturze klienta. Oznacza to, że dane urzędu czy instytucji nigdy nie opuszczają jej środowiska – nie trafiają do zewnętrznych serwerów, nie są przetwarzane przez podmioty trzecie. Dla administracji publicznej, odpowiedzialnej za dane obywateli, to nie jest opcja dodatkowa, lecz warunek konieczny, by w ogóle rozważać wdrożenie AI.

Przewidywalność to kolejny fundament naszych produktów. Instytucje publiczne potrzebują narzędzi, których działanie jest powtarzalne, można je audytować, a także koszt ich wdrożenia i działania jest stabilny – to buduje zaufanie zarówno wśród urzędów, jak i obywateli korzystających z usług publicznych. Stabilność wynika z faktu, że nasze systemy rozwijamy z myślą o długofalowej współpracy z sektorem, a nie o jednorazowym wdrożeniu.

Ważne jest też to, że Safe Stack jest polską firmą. Rozumiemy lokalny kontekst prawny, organizacyjny i kulturowy administracji publicznej w Polsce – wiemy, jak wygląda proces wdrożenia w urzędzie i jakie wymagania stawiają konkretne przepisy krajowe. To realna przewaga wobec zagranicznych dostawców, którzy rzadko oferują taki poziom zgodności i wsparcia dostosowanego do polskich realiów.

Widzę, że administracja publiczna coraz odważniej sięga po sztuczną inteligencję, ale robi to selektywnie, wybierając partnerów łączących innowacyjność z odpowiedzialnością. Jestem przekonany, że bezpieczne, sprawdzalne i zgodne prawnie rozwiązania działające lokalnie będą wyznaczać standard wdrożeń AI w sektorze publicznym w najbliższych latach.

Suwerenność technologiczna - wprowadzenie

W obliczu dynamicznych przemian na arenie międzynarodowej i błyskawicznego rozwoju nowych technologii, kwestia niezależności technologicznej nabiera fundamentalnego znaczenia dla przyszłości Polski i całego kontynentu europejskiego. Przez niezależność tę rozumiemy gotowość państwa i jego gospodarki do samodzielnego budowania, kontrolowania i implementowania kluczowych technologii - poczynając od sektora obronnego i bezpieczeństwa energetycznego, przez transformację cyfrową, aż po zarządzanie łańcuchami dostaw.

Suwerenność technologiczna nabiera również szczególnego znaczenia w kontekście rozwoju sztucznej inteligencji. Wykorzystanie AI coraz mocniej wpisuje się w realia polskiego biznesu. Według raportu Polcom „Barometr cyfrowej transformacji polskiego biznesu 2025-2026”, niemal dwie trzecie firm (65%) korzysta z AI w dziedzinie cyberbezpieczeństwa - przede wszystkim do identyfikowania zagrożeń i wykrywania nieprawidłowości w danych. Niewiele mniej, bo 63% przedsiębiorstw, sięga po algorytmy sztucznej inteligencji - w tym zaawansowane modele językowe - w celu usprawnienia zarządzania wiedzą i wsparcia procesów decyzyjnych. Z kolei ponad połowa organizacji (52%) wykorzystuje AI do automatyzacji codziennych operacji, zarówno w biurze, jak i na produkcji. Co istotne, 37% firm aktywnie przygląda się dostępnym rozwiązaniom opartym na AI i planuje ich wdrożenie w niedalekiej przyszłości.

Nie bez znaczenia jest fakt, że dynamiczny rozwój sztucznej inteligencji idzie w parze z rozbudową zaplecza technologicznego. Fundamentem tej transformacji stają się chmura obliczeniowa i nowoczesne centra danych. Z danych Polcom wynika, że 62% firm z sektora MŚP zamierza zainwestować w rozwiązania chmurowe do 2026 roku, postrzegając je jako niezbędny warunek dalszego rozwoju i skalowania projektów opartych na AI.

Komitet Przełomowych Technologii działający przy Polskiej Izbie Informatyki i Telekomunikacji opracował dokument, w którym określa kompleksowe podejście do wzmacniania technologicznej odporności kraju. Określono w nim również warunki niezbędne do tego, by Polska mogła sprawnie funkcjonować w szybko ewoluującym środowisku cyfrowym.

W dokumencie wyraźnie zaznaczono, że suwerenność technologiczna nie jest równoznaczna z dążeniem do samowystarczalności czy odcinaniem się od światowych ekosystemów technologicznych. Jej sedno tkwi w zdolności państwa do autonomicznego podejmowania decyzji, sprawowania kontroli nad danymi i infrastrukturą krytyczną oraz utrzymania ciągłości działania nawet w obliczu sytuacji kryzysowych⁰¹.

Kluczowe jest wypracowanie kompleksowej i długofalowej strategii, która wyznaczy kierunki rozwoju i połączy politykę państwa w obszarach cyfryzacji, nowoczesnych technologii oraz cyberbezpieczeństwa - bez podziałów resortowych i powielania tych samych działań przez kolejne ekipy rządzące.

Niezbędne jest również krajowe wsparcie branży technologicznej poprzez tworzenie przyjaznego środowiska dla polskich firm z tego sektora, oferowanie preferencyjnych instrumentów finansowych, pomoc w wejściu

na rynki zagraniczne, a także dostrzeganie wartości kapitału ludzkiego - aby utalentowani specjaliści decydowali się pozostać i pracować w Polsce, zamiast wyjeżdżać do wielkich międzynarodowych korporacji i budować ich potencjał oraz zyski.

Nieodzowne są również reformy systemu edukacji, które wprowadzą do podstawy programowej zagadnienia z zakresu cyberbezpieczeństwa i nowych technologii, takich jak sztuczna inteligencja czy technologie kwantowe, a także kształcenie umiejętności krytycznego myślenia. To ostatnie powinno zwiększyć odporność społeczeństwa na dezinformację, stanowiącą dziś nieodłączny element cyfrowej rzeczywistości. Wzrost świadomości technologicznej i większa odporność na manipulację informacyjną przełożą się na wzrost zaufania obywateli do państwa i jego instytucji, a także na chęć rozwijania kompetencji z dziedziny cyberbezpieczeństwa.

Nowoczesne podejście do nauczania powinno obejmować wszystkie szczeble edukacji, w tym wsparcie dla nauczycieli w systematycznym podnoszeniu kwalifikacji. Nieustanny rozwój wiedzy i umiejętności dotyczy również pracowników sektora publicznego i prywatnego - jak wynika z najnowszego raportu „The Future of Jobs 2025” opracowanego przez ekspertów Światowego Forum Ekonomicznego, zdolność do ciągłego uczenia się stanie się kluczową kompetencją w dzisiejszym świecie. Ten sam doku-

ment wskazuje na rosnące zapotrzebowanie na specjalistów ds. bezpieczeństwa i cyberbezpieczeństwa. Tym bardziej zatem kształcenie społeczeństwa w tym kierunku stanie się warunkiem koniecznym dla budowania silnego i nowoczesnego państwa.

Technologiczna suwerenność Polski powinna być budowana na fundamencie silnych partnerstw z krajami UE i NATO, obejmujących wymianę informacji o zagrożeniach, wspólne ćwiczenia oraz skoordynowane reagowanie na dynamicznie zmieniające się cyberzagrożenia. W tym zakresie aktywnie działają już Wojska Obrony Cyberprzestrzeni, które nawiązują sojusze, biorą udział w międzynarodowych manewrach i wymieniają doświadczenia z dowódcami jednostek cyberwojsk z krajów demokratycznych.

Do kolejnych czynników budujących technologiczną suwerenność Polski należy wspieranie badań i rozwoju w sektorach - które należałoby już określać mianem teraźniejszości, nie przyszłości - takich jak półprzewodniki, kryptografia, sztuczna inteligencja czy blockchain.

Cyfrowa przyszłość Polski musi być zakorzeniona w technologicznej suwerenności i cyberbezpieczeństwie, rozumianych jako integralne elementy spójnej, całościowej strategii⁰².

Suwerenna chmura

Koncepcja suwerennej chmury obliczeniowej (ang. sovereign cloud) wyłoniła się jako odpowiedź na rosnące napięcia między potrzebą korzystania z globalnej infrastruktury cyfrowej a koniecznością zachowania kontroli nad danymi i systemami o znaczeniu strategicznym. Jak podkreślają eksperci z Instytutu Łączności, w przypadku administracji publicznej nie jest to jedynie kwestia wyboru technologicznego, lecz wymóg niezbędny z perspektywy modernizacji państwa i budowy fundamentów technologicznej niezależności kraju⁰³. Z kolei przedstawiciele Centralnego Ośrodka Informatyki (COI) zaznaczają, że infrastruktura dla kluczowych systemów musi zapewniać ciągłość funkcjonowania państwa, a dane zawierające informacje niejawne w ogóle nie mogą być lokowane w publicznych środowiskach chmurowych⁰⁴.

Suwerenna chmura obliczeniowa różni się od tradycyjnych usług chmurowych przede wszystkim zakresem jurysdykcji i kontroli operacyjnej. Według ekspertów Związku Cyfrowa Polska dane wrażliwe powinny być przechowywane w suwerennych środowiskach, najlepiej zarządzanych przez krajowych dostawców, co gwarantuje pełną kontrolę. Zgodnie z wytycznymi samego COI, od dostawców komercyjnych współpracujących z państwem oczekuje się ograniczenia fizycznej lokalizacji obsługi danych do terytorium Polski. Ma to szczególne znaczenie w kontekście przepisów takich jak amerykańska ustawa CLOUD Act, która w określonych okolicznościach umożliwia organom ścigania USA dostęp do danych przechowywanych przez amerykańskich dostawców. Administracja publiczna czy instytucje odpowiedzialne za infrastrukturę krytyczną nie mogą akceptować takiego ryzyka.

Warto jednak podkreślić, że wyzwanie to dotyczy nie tylko szczebla centralnego. Według badań przeprowadzonych przez Związek Cyfrowa Polska i kancelarię Andersen Tax & Legal, to właśnie jednostki samorządu terytorialnego wykazują dziś największą otwartość na technologie chmurowe – korzysta z nich aż 79 proc. badanych jednostek. Gminy czy powiaty przetwarzają każdego dnia ogromne wolumeny danych wrażliwych obywateli.

Jak zauważają przedstawiciele Państwowego Instytutu Badawczego NASK, mniejsze instytucje często nie dysponują znacznymi nakładami finansowymi na zakup własnej infrastruktury IT. W praktyce, przy braku strategii nadzorczej, dane mieszkańców mogą trafiać do infrastruktury globalnych dostawców, co prowadzi do ryzykownego zjawiska uzależnienia od jednej firmy (vendor lock-in), przed którym wielokrotnie ostrzegają autorzy wspomnianego raportu pt. "Polska administracja publiczna w drodze do chmury". Skala tego problemu jest znaczna – każda spośród ponad 2 400 gmin i blisko 380 powiatów jest administratorem w rozumieniu RODO.

Suwerenna chmura na poziomie samorządowym oznacza w praktyce dostęp do bezpiecznej, wspólnotowej infrastruktury chmurowej. Zgodnie z oficjalnymi informacjami rządowymi, Rządowa Chmura Obliczeniowa (RChO) funkcjonuje już jako wspólnotowa chmura wyposażona w Rządowy Klaster Bezpieczeństwa (RKB) oraz Standardy Cyberbezpieczeństwa Chmur Obliczeniowych (SCCO)⁰⁵. Co ważne, jak zapowiadają przedstawiciele Departamentu Wspólnej Infrastruktury Informatycznej Państwa działającego w ramach COI, państwo pracuje obecnie nad koncepcją centralnie nadzorowanej chmury obliczeniowej właśnie dla samorządów. Dzięki temu to strona rządowa dbałaby o zgodność

postępowań przetargowych z rygorystycznymi wymogami bezpieczeństwa, zdejmując ten ciężar z urzędów gmin i pozwalając im bezpiecznie korzystać z nowoczesnych rozwiązań. Krok ten wpisuje się w postulowaną przez analityków rynku pilną potrzebę wdrożenia krajowej polityki Cloud First na wszystkich szczeblach administracji.

Rządowa Chmura Obliczeniowa – stan obecny i wyzwania

Kluczowym projektem w obszarze suwerennej chmury dla polskiej administracji jest Rządowa Chmura Obliczeniowa (RChO), realizowana przez Ministerstwo Cyfryzacji przy udziale Centralnego Ośrodka Informatyki. RChO to chmura wspólnotowa administracji publicznej – infrastruktura przeznaczona do wyłącznego użytku przez określoną grupę organizacji mających wspólne założenia, w tym misję, wymagania bezpieczeństwa i politykę zgodności z regulacjami⁰⁶. Chmura Rządowa funkcjonuje w ramach kompleksowego systemu zabezpieczeń o nazwie Rządowy Klaster Bezpieczeństwa, zaprojektowanego zgodnie z architekturą zero-trust oraz rekomendacjami zebranymi w Narodowych Standardach Cyberbezpieczeństwa (NSC), a dla jej potrzeb działają dedykowane centra Security Operations Center (SOC) oraz Network Operations Center (NOC)⁰⁷. Architektura ta ma zapewniać ciągłość i bezpieczeństwo działania systemów na szczeblu centralnym, przy czym trwają prace nad udostępnieniem centralnie nadzorowanej chmury obliczeniowej również dla jednostek samorządowych, które dotychczas korzystały z chmury głównie na zasadach komercyjnych.

Mimo tak zaawansowanych inwestycji, dostępne dane wskazują na istotne bariery wdrożeniowe krajowych rozwiązań. Wbrew obiegowym opiniom, to nie jednostki poniżej szczebla centralnego mają z tym największy problem. Przeciwnie – jednostki samorządu terytorialnego wykazują największą otwar-

tość na technologie chmurowe (dane w chmurze przechowuje 79% z nich), podczas gdy administracja rządowa na szczeblu centralnym wykazuje o wiele większą ostrożność.

Wyżej wspomniany raport pt. "Polska administracja publiczna w drodze do chmury" opracowany z udziałem ekspertów wskazuje, że niska adopcja RChO i systemu Zapewniania Usług Chmurowych (ZUCH) wynika m.in. z braku zaufania, nieczytelności instytucjonalnej, słabej oferty technicznej oraz braku wsparcia wdrożeniowego. Eksperti zaznaczają ponadto, że istotnymi barierami rozwoju są nieefektywny system finansowania oparty na rezerwie celowej oraz stosunkowo krótki czas obecności tych państwowych rozwiązań na rynku. W praktyce sektor publiczny wciąż nie traktuje ZUCH jako realnej alternatywy wobec globalnych dostawców⁰⁸.

Problem ten uderza w suwerenność technologiczną i bezpieczeństwo państwa: jednostki administracji nagminnie sięgają po rozwiązania globalne kierując się głównie dotychczasowymi umowami licencyjnymi. Powoduje to niebezpieczne zjawisko uzależnienia od jednego dostawcy (vendor lock-in) – aż 78% jednostek korzystających z chmury używa rozwiązań zaledwie jednej firmy (Microsoft). Co gorsza, instytucje publiczne dokonują wdrożeń punktowo, bez długofalowych strategii i zabezpieczeń kontraktowych przewidujących procedury wyjścia czy migracji danych. Obserwację rażąco niskiej popularności państwowych rozwiązań chmurowych potwierdziła Najwyższa Izba Kontroli w wystąpieniu pokontrolnym z marca 2025 roku, wykazując, że w latach 2020-2022 aktywność w systemie ZUCH podjęło zaledwie 9 podmiotów⁰⁹.

Kierunki działania

Strategiczne wyzwanie dla Polski polega na wypracowaniu spójnej polityki chmurowej działającej na wszystkich szczeblach admini-

stracji – nie tylko w KPRM czy resortach, ale też w starostwie powiatowym czy gminnym centrum usług wspólnych. Niezbędne jest określenie katalogu danych i systemów, które ze względu na swój charakter muszą być przetwarzane wyłącznie w infrastrukturze podlegającej pełnej krajowej lub europejskiej kontroli, i stopniowe przenoszenie tych zasobów do środowisk spełniających wymogi wynikające z RODO, NIS 2, DORA i ustawy o KSC. Dla pozostałych zastosowań możliwe jest korzystanie z rozwiązań hybrydowych, pod warunkiem właściwej klasyfikacji danych i odpowiednich zabezpieczeń kontraktowych. Kluczowe jest przy tym, by mniejsze jednost-

ki administracji nie były pozostawione z tym zadaniem same – bez wsparcia kompetencyjnego, bez gotowych wzorców umownych i bez dostępu do bezpiecznej infrastruktury wspólnotowej, która nie wymaga od nich własnych zespołów IT.

Suwerenna chmura to nie alternatywa dla cyfryzacji – to jej warunek konieczny wszędzie tam, gdzie dane obywateli są przetwarzane w imieniu państwa, niezależnie od tego, czy dzieje się to w Warszawie, czy w urzędzie gminy liczącej kilka tysięcy mieszkańców.



Lokalne modele AI

Dyskusja o suwerenności cyfrowej nie może ograniczać się wyłącznie do infrastruktury chmurowej. Równie istotne – a może nawet bardziej – jest pytanie o to, jakie modele sztucznej inteligencji przetwarzają dane pracowników administracji, urzędników, lekarzy, prawników i pracowników firm. Rzeczywistość jest taka, że AI weszła już do codziennej pracy i to głębiej, niż większość instytucji zdaje sobie sprawę. Według raportu KPMG z 2025 roku już 69% Polaków regularnie korzysta z narzędzi opartych na sztucznej inteligencji, co plasuje Polskę powyżej średniej światowej¹⁰. W samej administracji publicznej obraz jest jeszcze bardziej wymowny – badanie NASK „AI w e-administracji publicznej” przeprowadzone na grupie ponad 6 000 urzędników wykazało, że 77% z nich ma kontakt ze sztuczną inteligencją przynajmniej raz dziennie, a niemal połowa – 46% – aktywnie wykorzystywała narzędzia generatywnej AI do realizacji swoich obowiązków zawodowych w ciągu ostatnich sześciu miesięcy¹¹. Pytanie nie brzmi więc czy pracownicy używają AI, ale czy organizacje mają kontrolę nad tym, dokąd trafiają wpisywane przez nich dane.

Gdy pracownik wkleja fragment dokumentu do zewnętrznego modelu językowego dostępnego przez przeglądarkę, dane te opuszczają organizację. Trafiają na serwery dostawcy – najczęściej poza granicami Polski i Unii Europejskiej – gdzie mogą być przetwarzane, przechowywane i potencjalnie wykorzystane do dalszego trenowania modeli. W przypadku danych osobowych narusza to RODO. W przypadku danych objętych tajemnicą służbową, bankową czy lekarską – narusza przepisy branżowe. W przypadku informacji niejawnych – stanowi zagrożenie dla bezpieczeństwa państwa. Codzienne korzystanie z komercyjnych asystentów AI bez odpowiednich zabezpieczeń to jeden z najbardziej niedocenianych wektorów wycieku danych w nowoczesnej organizacji – tym bardziej niepokojący, że skala tego zjawiska jest już dziś masowa, a nie marginalna.

Odpowiedzią na ten problem są lokalne modele językowe – systemy sztucznej inteligencji uruchamiane i przetwarzające dane wyłącznie na infrastrukturze kontrolowanej przez daną organizację lub państwo. Dane nie opuszczają środowiska wewnętrznego, model działa bez połączenia z zewnętrznymi

serwerami, a organizacja ma pełną kontrolę nad tym, co model robi, jak działa i kto ma do niego dostęp.

Czym jest lokalny model językowy i jak go uruchomić

Duży model językowy (ang. Large Language Model, LLM) to system sztucznej inteligencji wytrenowany na ogromnych zbiorach tekstu, który potrafi rozumieć i generować język naturalny – odpowiadać na pytania, streszczać dokumenty, redagować pisma, tłumaczyć, analizować treść umów czy odpowiadać na zapytania obywateli. To ta sama technologia, która stoi za popularnymi asystentami AI, takimi jak ChatGPT czy Gemini. Różnica nie leży więc w tym, co model potrafi, a w tym, gdzie i przez kogo jest uruchamiany.

Kluczowa różnica między modelem lokalnym a usługą zewnętrzną leży w miejscu przetwarzania. Lokalny model językowy jest pobierany i uruchamiany bezpośrednio na serwerach organizacji – może to być serwer w centrum danych urzędu, instytucji publicznej, szpitala czy firmy. Nie wymaga stałego połączenia z internetem ani przekazywania danych do

dostawcy zewnętrznego. Administratorzy IT organizacji zarządzają całym środowiskiem: instalacją, aktualizacjami, kontrolą dostępu i logowaniem zapytań. To właśnie ta kontrola – a nie sama jakość modelu – jest wartością strategiczną z perspektywy bezpieczeństwa danych i zgodności z prawem.

Wdrożenie lokalnego modelu językowego – szczególnie w wariancie open source – jest dziś technicznie dostępne dla organizacji dysponujących podstawowym zapleczem serwerowym wraz z odpowiednimi kartami graficznymi zdolnymi do inferencji, dedykowanymi dla sztucznej inteligencji. Obecnie nie jest już to infrastruktura zarezerwowana wyłącznie dla dużych korporacji – jest to sprzęt dostępny w standardowych budżetach IT średniej wielkości urzędu czy szpitala powiatowego. Większe modele wymagają bardziej rozbudowanego zaplecza obliczeniowego, ale dla podmiotów publicznych możliwe jest korzystanie ze wspólnej infrastruktury w obrębie suwerennej chmury, lub serwerów z większą liczebnością kart graficznych.

Ważnym elementem lokalnego wdrożenia jest też możliwość tzw. dostrajania modelu (ang. fine-tuning) – czyli dalszego trenowania już gotowego modelu na danych specyficznych dla danej organizacji lub dziedziny. Urząd, który dysponuje własnymi zasobami dokumentów administracyjnych, aktów prawnych czy korespondencji urzędowej, może na ich podstawie poprawić jakość odpowiedzi modelu w swoim konkretnym kontekście. Dane użyte do dostrajania pozostają wewnątrz organizacji, a sam model staje się coraz lepiej dopasowany do realnych potrzeb – bez konieczności angażowania zewnętrznego dostawcy.

Modele open source jako fundament niezależności – przykład Mistral AI

Szczególną rolę w budowaniu europejskiej niezależności od modeli pozaeuropejskich

odgrywają modele o otwartych wagach (open-weight), których twórcą jest m.in. paryski Mistral AI. Historia tej firmy jest sama w sobie argumentem na rzecz tezy, że Europa nie musi być wyłącznie odbiorcą technologii AI tworzonych w Stanach Zjednoczonych czy Chinach.

Mistral AI to francuska firma założona w 2023 roku przez byłych badaczy z Google DeepMind i Meta¹². W ciągu zaledwie dwóch lat z niewielkiego startupu stała się podmiotem o globalnym zasięgu – jej wycena na koniec 2025 roku przekroczyła 14 miliardów dolarów¹³, co istotne, wzrost ten nie odbył się kosztem porzucenia filozofii otwartości – Mistral konsekwentnie udostępnia znaczną część swoich modeli publicznie, co odróżnia go od większości globalnych konkurentów. Filozofia Mistrala jest bezpośrednio zbieżna z postulatami suwerenności cyfrowej. Jednym z filarów strategii firmy jest wsparcie dla ruchu open source – modele Mistrala mogą być uruchamiane i rozwijane niezależnie, na własnej infrastrukturze, bez uzależnienia od jednej platformy dostępowej. To odpowiedź na zjawisko vendor lock-in, które szczególnie dotyka firmy technologiczne, instytucje badawcze i administrację publiczną korzystające z zamkniętych ekosystemów globalnych dostawców¹⁴.

Równie ważna jak sama technologia jest ścieżka, którą Mistral obrał w relacjach z sektorem publicznym. W lipcu 2025 roku firma uruchomiła inicjatywę „AI for Citizens”, skierowaną do państw i instytucji publicznych, której celem jest strategiczne wykorzystanie AI w transformacji usług publicznych. W ramach tej inicjatywy Mistral zawarł partnerstwa z armią francuską, francuską agencją pracy, rządem Luksemburga oraz różnymi europejskimi organizacjami sektora publicznego¹⁵.

Na poziomie infrastrukturalnym Mistral konsekwentnie realizuje wizję europejskiej suwerenności obliczeniowej – w marcu 2026 roku firma pozyskała 830 milionów dolarów

na budowę własnych centrów danych pod Paryżem i w Szwecji, dążąc do zapewnienia 200 MW mocy obliczeniowej w Europie do końca 2027 roku¹⁶.

Mistral jest zatem nie tylko dostawcą technologii, lecz aktywnie współkształtuje europejski model wdrażania AI – zarówno w sektorze publicznym, jak i prywatnym. Jego ścieżka rozwoju pokazuje, że niezależność od doliny krzemowej jest osiągalna, komercyjnie opłacalna i możliwa do zrealizowania w ramach europejskich wartości dotyczących ochrony danych i przejrzystości. Dla Polski i innych krajów Unii Europejskiej Mistral stanowi punkt odniesienia – zarówno jako dostawca modeli możliwych do wdrożenia lokalnie, jak i jako dowód, że europejski ekosystem AI może konkurować globalnie bez rezygnowania z kontroli nad danymi.

Bielik – polska odpowiedź na potrzeby językowe administracji

O ile Mistral oferuje modele wielojęzyczne o ogólnym przeznaczeniu, o tyle polskie instytucje posługują się językiem, który globalnym modelom sprawia realną trudność. Wysoki stopień złożoności polskiej fleksji, w połączeniu z hermetyczną terminologią prawną-administracyjną, stylem urzędowym oraz specyfiką aktów prawnych, stanowi barierę dla modeli trenowanych głównie na danych anglojęzycznych. W tych obszarach modele globalne wypadają nieporównywalnie słabiej niż rozwiązania zoptymalizowane pod rodzimą strukturę językową. Tę lukę wypełnia projekt Bielik¹⁷.

Bielik narodził się w 2024 roku z inicjatywy Fundacji SpeakLeash, przy wsparciu obliczeniowym Akademickiego Centrum Komputerowego Cyfronet AGH w Krakowie. Model był trenowany na polskich zbiorach danych – od dokumentów prawnych i ustaw, przez Projekt

Gutenberg, po polską Wikipedię – dzięki czemu nauczył się rozumieć niuanse języka, stylu i kultury, z którymi globalne modele mają problem¹⁸. Istotne jest też to, że dane użyte do trenowania były starannie wyselekcjonowane i udokumentowane pod kątem pochodzenia – co w kontekście zastosowań publicznych ma znaczenie zarówno prawne, jak i etyczne¹⁹. Bielik to projekt non-profit dostępny na licencji open source – można go pobrać i uruchomić lokalnie, bez konieczności korzystania z chmury, na CPU i zwykłych GPU. Właśnie ta właściwość jest kluczowa z perspektywy administracji: dane wpisywane przez urzędnika nigdzie nie trafiają, bo model działa w całości wewnątrz organizacji, pod kontrolą jej własnych administratorów.

Projekt zyskał szybko uznanie zarówno w kraju, jak i za granicą. W 2024 roku Fundacja SpeakLeash otrzymała wyróżnienie „AI Spotlight” podczas konferencji GOSIM 2024 oraz nagrodę specjalną Masters&Robots²⁰. W maju 2025 roku Bielik uzyskał honorowy patronat Ministerstwa Cyfryzacji i znalazł się w gronie dziesięciu najbardziej wpływowych projektów open source AI na świecie podczas konkursu Spotlight AI 2025²¹. Patronat ministerialny nie jest wyłącznie symboliczny – oznacza instytucjonalne uznanie Bielika za element krajowej strategii cyfrowej i otwiera drogę do jego wdrożeń w administracji publicznej.

Bielik, podobnie jak Mistral, pokazuje że lokalny model językowy to nie kompromis wobec globalnych rozwiązań, lecz świadoma decyzja o zachowaniu kontroli nad danymi przy jednoczesnym utrzymaniu wysokiej jakości narzędzia. W przypadku Bielika do tego argumentu dochodzi jeden dodatkowy: jest to model rozumiejący Polskę – jej język, prawo i realia administracyjne – czego zagraniczne modele nie są w stanie zapewnić w tym samym stopniu.

PLLuM – model finansowany ze środków publicznych, gotowy dla administracji

Obok Bielika, który jest projektem społecznościowym fundacji non-profit, w Polsce istnieje drugi kluczowy model językowy – PLLuM (Polish Large Language Model) – finansowany bezpośrednio ze środków publicznych i celowo zaprojektowany z myślą o zastosowaniach w administracji publicznej.

Konsorcjum PLLuM zostało zawiązane 29 listopada 2023 roku przez sześć wiodących polskich instytucji naukowych: Politechnikę Wrocławską jako lidera projektu, Państwowy Instytut Badawczy NASK, Ośrodek Przetwarzania Informacji, Instytut Podstaw Informatyki PAN, Uniwersytet Łódzki oraz Instytut Slawistyki PAN. Budżet pierwszego etapu projektu wyniósł 14,5 mln zł, a Ministerstwo Cyfryzacji przeznaczyło następnie kolejne 19 mln zł na wdrożenie modelu w administracji publicznej przez nowych partnerów – Centralny Ośrodek Informatyki i Cyfronet AGH²². PLLuM to rodzina modeli językowych zaprojektowana specjalnie z myślą o języku polskim – jego strukturze i stylach komunikacji – tak aby generować precyzyjne i zrozumiałe treści w języku polskim. Modele PLLuM zostały zabezpieczone przed generowaniem treści szkodliwych i toksycznych, co jest kluczowe przy zastosowaniu w administracji publicznej. Model operuje na strukturach od 8 do 70 miliardów parametrów, co pozwala dopasować jego rozmiar i wymagania sprzętowe do możliwości konkretnej instytucji. Zastosowania są konkretne i bezpośrednio odpowiadają na codzienne potrzeby urzędów. PLLuM ma dostarczyć wirtualnego asystenta w aplikacji mObywatel, który wesprze obywateli w uzyskiwaniu informacji publicznych, a także inteligentnego asystenta urzędniczego, który zautomatyzuje przetwarzanie dokumentów, analizę treści, wyszukiwanie informacji i wsparcie w udzielaniu odpowiedzi na pytania obywateli²³. Pierwsze wdro-

żenie pilotażowe uruchomiono w Urzędzie Miasta Częstochowy – to pierwsze tego typu porozumienie w Polsce, pozwalające przetestować praktyczne wykorzystanie polskiej sztucznej inteligencji w samorządzie, gdzie AI wspiera urzędników m.in. w tworzeniu pism urzędowych, streszczaniu dokumentów oraz analizie zapytań od obywateli. Gdynia z kolei zintegrowała PLLuM z wyszukiwarką Biuletynu Informacji Publicznej, umożliwiając mieszkańcom zadawanie pytań w języku naturalnym bezpośrednio w zasobach BIP²⁴. Projekt PLLuM przekształca się w projekt HIVE, zmierzający do budowy szerszego ekosystemu polskich modeli językowych. Liderem konsorcjum HIVE zostaje Ośrodek Badań nad Bezpieczeństwem SI w NASK, a Cyfronet AGH będzie dostarczać mocy obliczeniowych pod uczenie modeli i ich udostępnianie. Ministerstwo Cyfryzacji zapowiada, że PLLuM i Bielik mają ze sobą współpracować i wspólnie pozyskiwać dane, tworząc razem rodzinę polskich modeli językowych.

Lokalne modele a suwerenna chmura – spójność systemu

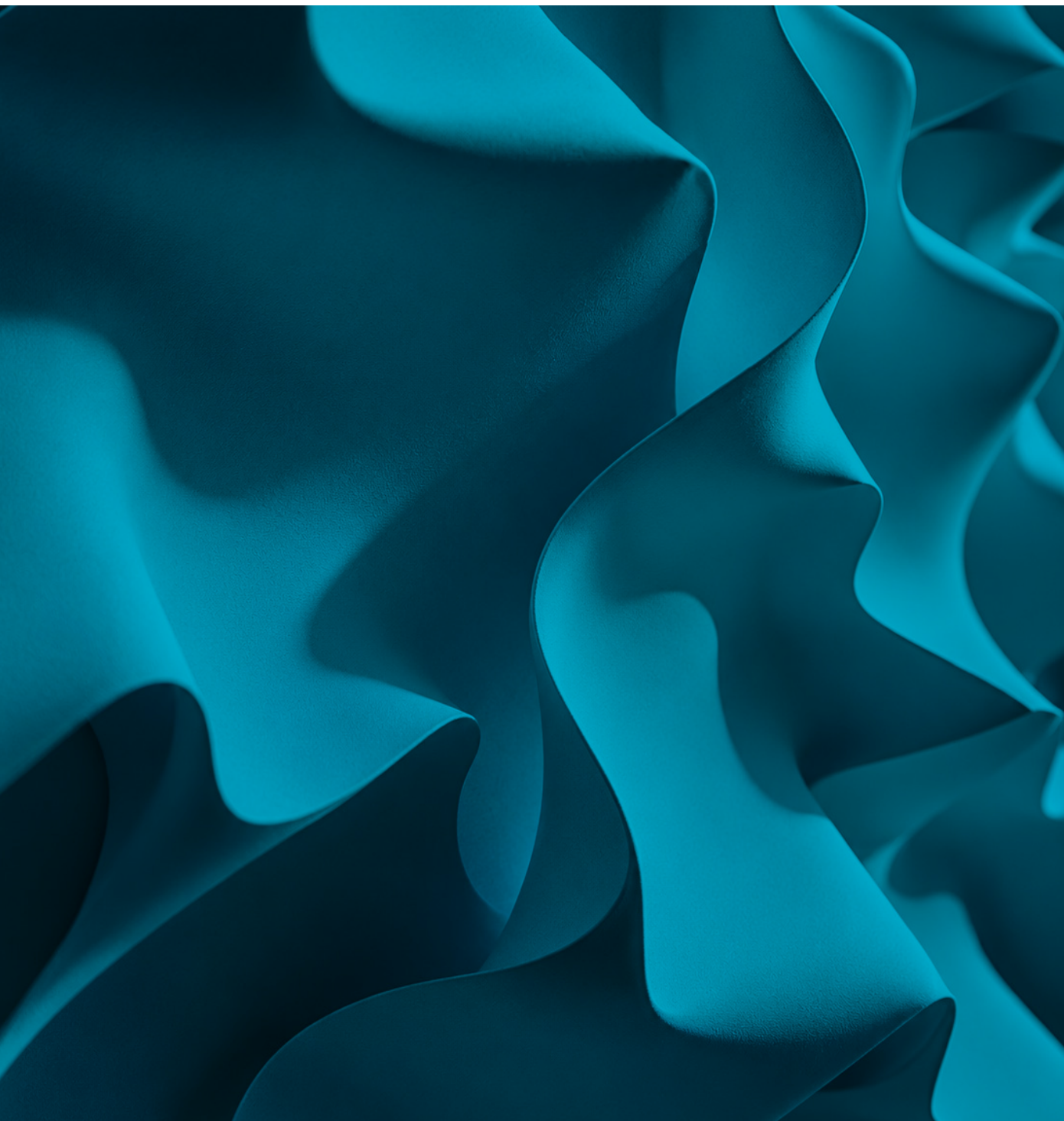
Lokalne modele językowe i suwerenna chmura obliczeniowa to dwa uzupełniające się filary tej samej strategii. Suwerenna chmura zapewnia bezpieczne środowisko przechowywania i przetwarzania danych – lokalny model językowy zapewnia, że narzędzie AI korzystające z tych danych nie wyprowadza ich poza to środowisko. Jedno bez drugiego jest niekompletne: można mieć dane w suwerennej chmurze, ale korzystać z zewnętrznego modelu AI, który te dane przetwarza poza kontrolą organizacji. Właśnie dlatego obie kwestie muszą działać w symbiozie, jako elementy jednej polityki cyfrowej.

PLLuM i Bielik wyróżniają się na tle innych modeli językowych pełną kontrolą nad modelem i brakiem zależności od zagranicznych dostawców, niskimi kosztami dzięki własnym rozwiązaniom oraz dostosowaniem do pol-

skich realiów, w tym terminologii administracji publicznej. To cechy, które w kontekście wymogów AI Act, RODO i ustawy o KSC mają bezpośrednie przełożenie na zdolność organizacji do wykazania zgodności z prawem – co przy zewnętrznym, zamkniętym modelu jest często niemożliwe lub bardzo utrudnione.

czy organizacje zapewnią im narzędzia, które nie stwarzają ryzyka prawnego i bezpieczeństwa. Lokalne modele językowe – Bielik, PL-LuM, Mistral wdrożony on-premise – są gotową odpowiedzią na tę potrzebę, dostępną już dziś, nie w perspektywie wieloletnich programów badawczych.

Kierunek jest jasny: pracownicy administracji publicznej i sektorów regulowanych i tak korzystają już z narzędzi AI. Pytanie nie dotyczy tego, czy im na to pozwolić, lecz tego,



Infrastruktura państwowa

Suwerenna chmura i lokalne modele językowe to bardzo ważny aspekt cyfryzacji i bezpieczeństwa kraju w dobie rewolucji cyfrowej spowodowanej dynamicznym rozwojem sztucznej inteligencji – ich wdrożenie jest możliwe wyłącznie wtedy, gdy istnieje ku temu odpowiednia infrastruktura obliczeniowa. Polska dysponuje już dziś solidnymi fundamentami w tym obszarze, a realizowane i planowane inwestycje wskazują, że kraj może świadomie zbudować pozycję regionalnego lidera infrastruktury AI.

Kluczowym ośrodkiem krajowej infrastruktury obliczeniowej pozostaje Akademickie Centrum Komputerowe Cyfronet AGH w Krakowie – najdłużej działające centrum superkomputerowe w Polsce, którego zasoby są dostępne dla nauki i gospodarki w ramach sieci PLGrid. ACK Cyfronet AGH dysponuje trzema centrami danych, własną siecią światłowodową oraz zapleczem technicznym działającym całą dobę przez 365 dni w roku, a jego flagowy superkomputer Athena osiąga moc obliczeniową ponad 7,7 PetaFlopsów, w tym prawie 240 PetaFlopsów dedykowanych obliczeniom AI²⁵. To właśnie superkomputery Cyfronetu – Helios i Athena – były fundamentem, na którym wytrenowano model językowy Bielik.

W wymiarze komercyjnym istotną rolę odgrywa Fabryka AI uruchomiona przez Beyond.pl w Poznaniu w maju 2025 roku. To pierwsza w regionie Europy Środkowej i Wschodniej komercyjna platforma suwerennej AI, oferująca GPU as a Service i AI as a Service, zlokalizowana w kampusie o zakontraktowanej mocy do 100 MW. Obie inicjatywy – akademicka i komercyjna – uzupełniają się wzajemnie: Cyfronet obsługuje potrzeby nauki i administracji publicznej, Beyond.pl dostarcza suwerennej mocy obliczeniowej dla firm i instytucji, które nie chcą korzystać z serwerów poza granicami kraju²⁶.

Najambitniejszym projektem infrastrukturalnym pozostaje jednak Baltic AI GigaFactory. W czerwcu 2025 roku Polska, jako lider konsorcjum obejmującego Litwę, Łotwę i Estonię, złożyła do Komisji Europejskiej wnioski o dofinansowanie z funduszu InvestAI budowy gigafabryki AI o szacowanej wartości 3 miliardów euro – największej tego typu inwestycji w Europie Środkowej i Wschodniej. Wśród partnerów projektu znalazły się Allegro, Orange Polska, CloudFerro, NASK, IDEAS NCBR oraz Cyfronet²⁷. Planowana infrastruktura ma być w całości zasilana energią odnawialną²⁸. Projekt zakłada wprost wspieranie rozwoju lokalnych modeli językowych, w tym PLLuM i Bielika, a także integrację generatywnej AI z narzędziami administracyjnymi i procesami urzędowymi.

Polska infrastruktura obliczeniowa ma realną szansę stać się spójnym ekosystemem – od akademickich centrów superkomputerowych, przez komercyjne fabryki AI, po infrastrukturę europejskiej skali – jeśli planowane inwestycje zostaną konsekwentnie zrealizowane. Jej dalsza rozbudowa to warunek, który mógłby jednocześnie odblokować rozwój krajowych modeli językowych w środowiskach naukowych, umożliwić administracji publicznej wdrażanie AI na infrastrukturze spełniającej wymogi RODO, NIS 2 i AI Act, oraz otworzyć sektorowi prywatnemu dostęp do suwerennej mocy obliczeniowej bez konieczności wywłaszczania danych poza granice kraju.

Inwestycja w infrastrukturę państwową to nie wydatek na sprzęt – to inwestycja w zdolność do samostanowienia technologicznego. Kraj dysponujący własną infrastrukturą obliczeniową może wdrażać AI na własnych warunkach, zgodnie z prawem, które sam

stosuje, i z kontrolą nad danymi, której żaden zewnętrzny dostawca nie jest w stanie zagwarantować w tym samym stopniu.



Zależność od globalnych dostawców i ryzyko outso- urcingu AI

Rynek sztucznej inteligencji rozwija się dziś w bezprecedensowym tempie i coraz wyraźniej staje się jednym z fundamentów globalnej gospodarki cyfrowej. Zgodnie z prognozami UNCTAD przedstawionymi w raporcie „Technology and Innovation Report 2025: Inclusive Artificial Intelligence for Development”, wartość tego sektora może sięgnąć nawet 4,8 biliona dolarów do 2033 roku. To poziom porównywalny z gospodarkami największych państw, co pokazuje, że AI przestała być jedynie wizją przyszłości, a stała się szybko rosnącą branżą realnie wpływającą na rynki, zatrudnienie i układ sił na świecie.

Tak dynamiczny rozwój rodzi jednak istotne wątpliwości dotyczące tego, kto faktycznie skorzysta na tej transformacji. Raport wskazuje, że sektor AI cechuje się dużą koncentracją kapitału oraz innowacji w rękach niewielkiej liczby podmiotów, głównie z krajów wysoko rozwiniętych, takich jak Stany Zjednoczone i Chiny. Może to prowadzić do pogłębiania nierówności ekonomicznych, ograniczenia konkurencji oraz zawłaszczania wiedzy technologicznej. W efekcie wiele państw rozwijających się może zostać zepchniętych na margines cyfrowej rewolucji, mając ograniczony dostęp do korzyści płynących z rozwoju AI.

W tej sytuacji warto analizować nie tylko potencjał gospodarczy sztucznej inteligencji, ale również jej skutki społeczne i rozwojowe. Kluczowe pozostaje pytanie, w jaki sposób wykorzystać rosnącą wartość tego rynku, aby nie pogłębiała globalnych podziałów, lecz wspierała zrównoważony rozwój i bardziej sprawiedliwy dostęp do innowacji.

Według raportu UNCTAD w rozwoju technologii AI prym wiodą przede wszystkim korporacje ze Stanów Zjednoczonych i Chin – to one przeznaczają największe budżety na badania, zarządzają kluczowymi platformami obliczeniowymi i gromadzą rozległe zbiory danych niezbędne do budowania modeli AI.

Taka koncentracja władzy rynkowej rodzi zjawisko zwane „oligopolem technologicznym”. Nowym graczom niezwykle trudno przebić się na ten rynek – barierą są zarówno gigantyczne koszty infrastruktury obliczeniowej, jak i ograniczony dostęp do dużych zasobów danych oraz trudności z pozyskaniem kapitału. W efekcie innowacje przestają powstawać w rozproszonych ośrodkach naukowych czy różnych gospodarkach świata, lecz skupiają się w rękach kilku korporacji i państw.

Skutki tego stanu rzeczy są złożone i wielopłaszczyznowe. Błyskawiczny postęp technologiczny, który napędzają największe podmioty AI, bez wątpienia przyspiesza wdrażanie przełomowych rozwiązań – generatywnych modeli językowych, systemów predykcyjnych w przemyśle czy medycynie. Jednocześnie jednak rośnie groźba monopolizacji branży oraz uzależnienia pozostałych krajów od dostawców technologii narzucających własne warunki – zarówno cenowe, jak i dotyczące dostępu do innowacji.

Dla państw rozwijających się sytuacja ta stanowi szczególnie poważne wyzwanie. Nie mogąc dorównać zasobom globalnych liderów, lokalne firmy skazane są głównie na wdrażanie gotowych produktów, zamiast aktywnie uczestniczyć w ich tworzeniu. Taka zależność podkopuje zdolność do samodzielnego kształtowania polityki technologicznej

i gospodarczej, a w dłuższej perspektywie utrwala podział świata na ośrodki wytwarzające innowacje i peryferia, które jedynie z nich korzystają.

Aby sztuczna inteligencja nie stała się kolejnym źródłem pogłębiającego się podziału między bogatymi a biednymi, trzeba aktywnie dążyć do sprawiedliwszego rozdziału płynących z niej korzyści. Pomocne mogą być tu przepisy antymonopolowe, które ograniczą przewagę największych graczy w branży technologicznej – zmuszą ich do większej otwartości i umożliwią mniejszym firmom dostęp do kluczowych zasobów, takich jak dane czy moc obliczeniowa.

Niemniej ważne jest tworzenie sprzyjającego gruntu dla lokalnej przedsiębiorczości i innowacji. Finansowanie startupów, małych firm oraz instytutów badawczych umożliwia opracowywanie rozwiązań skrojonych pod potrzeby konkretnych społeczeństw. Dzięki temu AI może być realnie użyteczna tam, gdzie naprawdę jest potrzebna – w rolnictwie, służbie zdrowia czy gospodarce zasobami naturalnymi.

Ważną rolę mogą odegrać też modele otwartego tworzenia technologii. Udostępnianie publicznych zbiorów danych, bibliotek open source i wspólnych platform naukowych przyspiesza przepływ wiedzy i obniża bariery wejścia dla podmiotów bez dużego kapitału. Otwartość w podejściu do AI nie tylko wyrównuje szanse na rynku, ale też ułatwia krajom rozwijającym się szybsze przyswajanie nowych technologii.

Warto również stawiać na współpracę sektora publicznego z prywatnym. Wspólne inwestycje w infrastrukturę cyfrową, centra przetwarzania danych i programy kształcenia mogą podnieść kompetencje technologiczne regionów dziś wykluczonych z cyfrowej transformacji. Kluczowe jest przy tym ak-

tywne uczestnictwo organizacji międzynarodowych, które dysponując funduszami i programami pomocowymi, mogą skutecznie wspierać budowanie globalnej równowagi w ekosystemie AI.

Jeśli te działania zostaną podjęte w sposób przemyślany i wzajemnie uzgodniony, mają realną szansę przełamać monopol wąskiej grupy potęg technologicznych i otworzyć przestrzeń dla bardziej różnorodnego świata innowacji. Co najważniejsze – stwarzają możliwość, by sztuczna inteligencja przestała być przywilejem najzamożniejszych, a stała się powszechnym narzędziem rozwoju społecznego i gospodarczego.

Szczególnie ważne jest opracowanie takich rozwiązań, które zapewnią bardziej sprawiedliwy i otwarty rozwój tej technologii. Na poziomie krajowym konieczne są inwestycje w edukację, rozwój lokalnych innowacji oraz infrastrukturę cyfrową, aby umożliwić budowanie własnych kompetencji i uniezależnienie się technologiczne. Z kolei na poziomie międzynarodowym istotną rolę odgrywają wspólne normy etyczne i regulacyjne, otwarte podejście do innowacji oraz mechanizmy finansowe wspierające państwa rozwijające się.

Jeżeli działania te będą ze sobą odpowiednio skoordynowane, sztuczna inteligencja może stać się narzędziem nie tylko napędzającym wzrost gospodarczy, ale również wspierającym realizację celów zrównoważonego rozwoju. W przeciwnym razie rynek AI może pogłębić istniejące nierówności, dzieląc świat na tych, którzy korzystają z jej rozwoju, i tych, którzy pozostają na marginesie cyfrowej transformacji. Kierunek, w jakim podąży świat, zależy od decyzji podejmowanych już dziś – zarówno przez rządy, jak i instytucje międzynarodowe oraz sektor prywatny²⁹.

Rola krajowych technologii

Ministerstwo Cyfryzacji przedstawiło nową, zaktualizowaną wersję dokumentu „Polityka rozwoju sztucznej inteligencji w Polsce do 2030 roku”, w której uwzględniono opinie i sugestie zgłoszone w trakcie konsultacji społecznych. Głównym założeniem tej strategii jest wyznaczenie przemyślanego i spójnego kierunku rozwoju AI w Polsce na najbliższe lata.

Dokument zakłada budowę zintegrowanego ekosystemu sztucznej inteligencji, który będzie sprawnie funkcjonował i wspierał zarówno gospodarkę, jak i społeczeństwo. Kluczowe znaczenie mają tu nowoczesna infrastruktura technologiczna oraz wysoko wykwalifikowani specjaliści. Wskazano m.in. na wykorzystanie takich zasobów jak sieci PLGrid, Fabryki AI w Poznaniu i Krakowie czy planowaną Baltic AI GigaFactory, podkreślając ich znaczenie dla nauki i biznesu.

Po analizie postulatów z konsultacji zdecydowano także o rozszerzeniu grona głównych ośrodków rozwijających AI. Od 1 stycznia 2027 r. dołączą do nich Łódź oraz Górnośląsko-Zagłębiowska Metropolia, które wyróżniają się dużym potencjałem rozwojowym.

W części poświęconej sprawnemu funkcjonowaniu państwa zaplanowano m.in. uruchomienie platformy AI HUB Poland, która ma wspierać zarządzanie oraz wdrażanie rozwiązań AI w sektorze publicznym. Zakłada się, że do 2030 roku większość usług administracyjnych będzie dostępna z wykorzystaniem sztucznej inteligencji.

Przewidziano także wsparcie dla samorządów w zakresie zarządzania procesami, np. poprzez pilotażowe wykorzystanie modelu językowego PLLuM.

W dokumencie wskazano najważniejsze obszary gospodarki, w których AI ma największy potencjał, nawiązując przy tym do inicja-

tyw Komisji Europejskiej. Dla każdego z nich mają powstać szczegółowe plany wdrożeń przygotowywane wspólnie z odpowiednimi ministerstwami. Podkreślono, że skuteczne wykorzystanie AI wymaga współpracy między różnymi sektorami, dostępu do danych, odpowiednich regulacji oraz wykwalifikowanych kadr.

Zwrócono również uwagę na potrzebę bliższej współpracy nauki z biznesem, aby szybciej przekładać wyniki badań na praktyczne rozwiązania. Planowane działania są zgodne z europejskimi strategiami dotyczącymi wdrażania AI i jej wykorzystania w nauce.

W zakresie wsparcia dla małych i średnich przedsiębiorstw oraz startupów przewidziano różne narzędzia, takie jak program „Cyfrowa wyprawka”, dostęp do mocy obliczeniowych czy możliwość testowania rozwiązań w tzw. piaskownicach regulacyjnych.

Jednym z głównych celów strategii jest rozwój sztucznej inteligencji w sposób bezpieczny i godny zaufania, służący poprawie jakości życia. Dokument zakłada m.in. stworzenie narzędzi do walki z dezinformacją, szczególnie tą generowaną przez AI.

Promowane są także rozwiązania wspierające dzieci, osoby starsze, osoby z niepełnościami oraz wykluczone społecznie. Jednocześnie podkreślono znaczenie ochrony danych osobowych jako kluczowego elementu wszystkich działań związanych z AI.

W zaktualizowanej wersji doprecyzowano kwestie finansowania, monitorowania i koordynacji działań, a także uporządkowano wskaźniki oceny postępów.

Cała strategia stanowi kompleksową odpowiedź na wyzwania związane z szybkim rozwojem technologii AI. Wyznacza konkretne kierunki działań państwa, łącząc potrzeby

administracji, nauki, biznesu i społeczeństwa. Jej realizacja ma przebiegać etapami i we współpracy z różnymi interesariuszami, tak aby zapewnić bezpieczne i odpowiedzialne wykorzystanie sztucznej inteligencji zgodne z europejskimi standardami. Ministerstwo zapowiada dalsze działania wdrożeniowe oraz regularne informowanie o postępach³⁰.



Wnioski strategiczne

Skala inwestycji, jaką globalni liderzy przeznaczają na rozwój sztucznej inteligencji, powinna być punktem wyjścia do każdej strategicznej rozmowy o polskiej i europejskiej suwerenności cyfrowej. W marcu 2026 roku OpenAI zamknęło rundę finansowania o wartości 122 miliardów dolarów przy wycenie spółki na poziomie 852 miliardów dolarów, z udziałem Amazon, NVIDIA, SoftBank i Microsoft. Anthropic z kolei pozyskało od samego Amazona łącznie 33 miliardy dolarów zobowiązań inwestycyjnych, przy bieżącej wycenie firmy sięgającej 965 miliardów dolarów. Całkowite wydatki OpenAI na infrastrukturę obliczeniową mogą sięgnąć kilkuset miliardów dolarów w ciągu trzech lat – to poziom, który jeszcze niedawno był zarezerwowany dla państw, a nie pojedynczych firm technologicznych³¹. Te liczby nie służą tu jako wyraz podziwu – służą jako punkt odniesienia dla skali problemu. Polska i Europa nie są w stanie konkurować z pojedynczymi firmami prywatnymi, które dysponują budżetami inwestycyjnymi porównywalnymi z PKB średnich państw. Nie taka jest też stawka. Stawką jest zdolność do tego, by nie być wyłącznie konsumentem technologii tworzonej gdzie indziej, lecz tworzyć ją na własnych warunkach dywersyfikując się od zachodniej lub wschodniej technologii. Suwerenność cyfrowa nie jest jednym projektem ani jedną decyzją – **jest wynikiem symbiozy trzech równoległych inwestycji: w infrastrukturę obliczeniową, w krajowe modele AI i w kadry** zdolne to wszystko budować, wdrażać i utrzymywać. Każdy z tych elementów z osobna jest niewystarczający. Infrastruktura

bez modeli to pusta serwerownia. Modele bez infrastruktury to projekt, który nie ma gdzie działać. Kadry bez jednego i drugiego to potencjał, który odpłynie tam, gdzie będzie mógł być wykorzystany np. w zagranicznym sektorze prywatnym. Konsekwencją tego rozumowania jest konieczność istotnego zwiększenia publicznych nakładów na rozwój krajowego ekosystemu AI – zarówno bezpośrednio, jak i poprzez instrumenty współfinansowania z sektorem prywatnym i funduszami europejskimi. Budżet projektu PLLuM wynoszący łącznie 33,5 mln zł, przy całym swoim znaczeniu symbolicznym i praktycznym, jest nieporównywalny z nakładami, jakie w tym samym czasie przeznaczają na AI pojedyncze firmy z Doliny Krzemowej. Oznacza to, że Polska musi działać mądrze tam, gdzie nie może działać tak samo dużo i kosztownie – koncentrując zasoby na obszarach, w których krajowa specyfika językowa, prawna i administracyjna daje przewagę, której globalny dostawca nie jest w stanie replikować. **Suwerenna chmura, lokalne modele językowe i krajowa infrastruktura** obliczeniowa to nie trzy osobne tematy – to trzy warstwy tej samej architektury bezpieczeństwa cyfrowego państwa. Państwa, które tej architektury nie zbudują, będą zmuszone korzystać z cudzej – na warunkach, które mogą się zmienić, w jurysdykcji, której nie kontrolują, z dostępem, który może zostać ograniczony. Historia ostatnich lat pokazuje, że uzależnienie technologiczne jest równie poważnym ryzykiem strategicznym jak uzależnienie energetyczne. Polska doświadczyła tego drugiego. Nie powinna czekać, aż doświadczy pierwszego.

Rozdział 6

Suwerenność technologiczna

1. ITwiz, Jaka powinna być strategia Polski rozwoju AI na lata 2026–2030?, <https://itwiz.pl/jaka-powinna-by-c-strategia-polski-rozwoju-ai-na-lata-2026-2030/>, dostęp: 25.03.2026.
2. Bezpieczeństwo Polskie, Suwerenność technologiczna Polski wobec wyzwań cyberbezpieczeństwa, <https://bezpieczenstwopolskie.pl/artykuly/suwerennosc-technologiczna-polski-wobec-wyzwan-cyberbezpieczenstwa/>, dostęp: 25.03.2026.

Suwerenna chmura

1. Fundacja Cyfrowa Polska, Chmura w administracji publicznej – raport, https://cyfrowapolska.org/wp-content/uploads/2025/06/Raport-Cyfrowa-Polska_chmura-w-administracji.pdf, dostęp: 25.03.2026.2.
2. ITwiz, Chmura rządowa początkiem szerszego użycia chmury w administracji publicznej, <https://itwiz.pl/chmura-rzadowa-zaczatkiem-szerszego-uzycia-chmury-w-administracji-publicznej/>, dostęp: 25.03.2026.
3. Chmura.gov.pl, Przewodnik dla odbiorcy usług RChO, <https://chmura.gov.pl/informacje/przewodnik-dla-odbiorcy-uslug-rcho>, dostęp: 25.03.2026.
4. Chmura.gov.pl, Rządowa Chmura Obliczeniowa, <https://chmura.gov.pl/informacje/rzadowa-chmura-obliczeniowa>, dostęp: 25.03.2026.
5. Fundacja Cyfrowa Polska, Rząd i samorządy coraz śmielej korzystają z chmury, ale wciąż bez spójnej strategii i odpowiednich zabezpieczeń, <https://cyfrowapolska.org/raport-rzad-i-samorzady-coraz-smielej-korzystaja-z-chmury-ale-wciaz-bez-spojnej-strategii-i-odpowiednich-zabezpieczen/>, dostęp: 25.03.2026.

Lokalne modele AI

1. KPMG, Sztuczna inteligencja w Polsce, <https://kpmg.com/pl/pl/wiedza/technologia/sztuczna-inteligencja-w-polsce.html>, dostęp: 25.03.2026.
2. AI.gov.pl, AI w e-administracji – perspektywa urzędnika, <https://ai.gov.pl/aktualnosci/AI-w-e-administracji-perspektywa-urzednika>, dostęp: 25.03.2026.
3. Wikipedia, Mistral AI, https://pl.wikipedia.org/wiki/Mistral_AI, dostęp: 25.03.2026.
4. Wikipedia, Mistral AI, https://en.wikipedia.org/wiki/Mistral_AI, dostęp: 25.03.2026.
5. MarketingOnline, Co oferuje Mistral AI – sztuczna inteligencja rodem z Paryża, <https://www.marketingonline.pl/blog-co-oferuje-mistral-ai-sztuczna-inteligencja-rodem-z-paryza/>, dostęp: 25.03.2026.
6. VentureBeat, Mistral launches Mistral 3 – a family of open models, <https://venturebeat.com/ai/mistral-launches-mistral-3-a-family-of-open-models-designed-to-run-on>, dostęp: 25.03.2026.
7. Intuition Labs, Mistral Large 3 MoE LLM explained, <https://intuitionlabs.ai/articles/mistral-large-3-moe-llm-explained>, dostęp: 25.03.2026.
8. Bielik AI, <https://bielik.ai/>, dostęp: 25.03.2026.
9. TJSoft, Bielik AI – polski model językowy, <https://tjsoft.pl/bielik-ai-polski-model-jezykowy/>, dostęp: 25.03.2026.
10. AGH, Bielik – polski model językowy powstał w AGH, <https://www.agh.edu.pl/aktualnosci/detail/bielik-polski-model-jezykowy-powstal-w-agh>, dostęp: 25.03.2026.

Rozdział 6

11. Wikipedia, Bielik (model językowy), [https://pl.wikipedia.org/wiki/Bielik_\(model_j%C4%99zykowy\)](https://pl.wikipedia.org/wiki/Bielik_(model_j%C4%99zykowy)), dostęp: 25.03.2026.
12. Mikrokontroler.pl, Polski model językowy Bielik AI – premiera i uznanie, <https://mikrokontroler.pl/2025/05/13/polski-model-jezykowy-bielik-ai-premiera-w-paryzu-i-globalne-uznanie/>, dostęp: 25.03.2026.
13. CRN, Powstał PLLuM – polski model językowy AI, <https://crn.pl/aktualnosci/powstal-pllum-polski-model-jezykowy-ai-pluum/>, dostęp: 25.03.2026.
14. ITwiz, Polski rząd udostępnia model językowy PLLuM, <https://itwiz.pl/polski-rzad-udostepnia-polski-model-jezykowy-pllum/>, dostęp: 25.03.2026.

Infrastruktura państwowa

1. Gov.pl, Polski model językowy PLLuM trafi do samorządu, <https://www.gov.pl/web/cyfryzacja/polski-model-jezykowy-pllum-trafi-do-samorzadu>, dostęp: 25.03.2026.
 2. Cyfronet, What we do, <https://www.cyfronet.pl/pl/about-us/what-we-do>, dostęp: 25.03.2026.
- CRN, Beyond.pl uruchomił fabrykę AI, <https://crn.pl/aktualnosci/beyond-pl-uruchomil-fabryke-ai/>, dostęp: 25.03.2026.
3. LBKP, <https://lbkp.pl/en/6537-2/>, dostęp: 25.03.2026.
 4. AI Business, Poland wants to be the heart of AI in Europe, <https://aibusiness.pl/en/Poland-wants-to-be-the-heart-of-AI-in-Europe-Baltic-AI-gigafactory-is-getting-closer-to-implementation/>, dostęp: 25.03.2026.

Zależność od globalnych dostawców i ryzyko outsourcingu AI

1. PARP, <https://www.parp.gov.pl/component/content/article/89386>, dostęp: 25.03.2026.

Rola krajowych technologii

1. Gov.pl, Polityka rozwoju sztucznej inteligencji w Polsce do 2030 roku, <https://www.gov.pl/web/cyfryzacja/polityka-rozwoju-sztucznej-inteligencji-w-polsce-do-2030-roku--bezpieczne-i-odpowiedzialne-wykorzystanie-sztucznej-inteligencji>, dostęp: 25.03.2026.

Wnioski strategiczne

1. Bithub, 122 miliardy dolarów – OpenAI rozbija bank Doliny Krzemowej, <https://bithub.pl/inwestycje/nowe-technologie/122-miliardy-dolarow-w-jeden-strzal-openai-rozbija-bank-doliny-krzemowej/>, dostęp: 25.03.2026.

08

Rekomendacje



krajowa
izba
gospodarki
cyfrowej

Rekomendacje i wnioski z raportu – wersja rozszerzona

Niniejszy rozdział stanowi pogłębione rozwinięcie wniosków zaprezentowanych w Executive Summary. Jego celem nie jest powtórzenie pojedynczych obserwacji z poszczególnych części raportu, lecz ułożenie ich w spójny, operacyjny plan działania dla administracji publicznej i jednostek samorządu terytorialnego. Rekomendacje zostały pogrupowane w pięć obszarów priorytetowych, które razem tworzą szkielet odpowiedzialnej i bezpiecznej transformacji opartej na sztucznej inteligencji.

Punkt wyjścia jest jednoznaczny: transformacja oparta na AI nie jest perspektywą – ona już się rozpoczęła, w dużej mierze oddolnie i poza ramami systemowymi. 46% pracowników administracji publicznej wykorzystuje narzędzia generatywnej AI w codziennej pracy, jednak aż 90% z nich sięga po rozwiązania publiczne, często darmowe, a 70% jako główną barierę wskazuje brak procedur. Pytanie nie brzmi już zatem „czy wdrażać AI”, lecz „jak zrobić to bezpiecznie, zgodnie z regulacjami i z realną korzyścią dla obywatela”. Poniższe rekomendacje są odpowiedzią na to pytanie.

1. Wdrożenie operacyjne – od shadow AI do bezpiecznego standardu pracy

Pierwszym i najpilniejszym obszarem działań jest uporządkowanie tego, co w administracji publicznej dzieje się już dziś – w sposób żywiołowy i często niekontrolowany. Skala oddolnego korzystania z generatywnej AI jest na tyle duża, że dalsze ignorowanie tego zjawiska lub próby jego administracyjnego ograniczenia oznaczałyby utratę realnej szansy na poprawę jakości pracy urzędów, przy jednoczesnym utrwaleniu istniejących ryzyk. Dane zebrane w raporcie pokazują, że 81% urzędników widzi potencjał AI w automatyzacji rutynowych zadań, 76% wskazuje oszczędność czasu, 77% docenia jej znaczenie w tłumaczeniach i komunikacji, a 72% – w pracy z dokumentami. Jednocześnie 67% obywateli opowiada się za szerszym stosowaniem AI w administracji.

Oznacza to, że istnieje silne zapotrzebowanie zarówno po stronie pracowników administracji, jak i obywateli, na rozwiązania, które technologia może zaadresować już dziś. Nie chodzi o zastępowanie urzędników, lecz o przejmowanie powtarzalnych, czasochłonnych czynności, które obciążają administrację i wydłużają obsługę spraw. Najwyższy potencjał wartości znajduje się w obszarach analizy i streszczania dokumentów, automatycznego przygotowywania odpowiedzi, wyszukiwania informacji w aktach, wielojęzycznej obsługi obywateli oraz wsparcia w pracy działów IT – od utrzymania systemów po tworzenie oprogramowania, gdzie obserwowane są przyrosty wydajności rzędu kilkudziesięciu procent.

- **Rekomendacja:** Administracja publiczna powinna porzucić logikę zakazów i przejść do logiki kontrolowanego umożliwiania. Zadaniem instytucji nie jest ograniczanie dostępu do technologii, z której pracownicy i tak korzystają, lecz wytypowanie 3–5 obszarów priorytetowych o największym potencjale wartości, uruchomienie w nich pilotaży opartych na zaufanych narzędziach, a następnie skalowanie sprawdzonych rozwiązań w sposób zintegrowany z istniejącymi systemami. Każde

wdrożenie powinno mieć jasno zdefiniowany cel operacyjny, mierzalne wskaźniki sukcesu oraz właściciela po stronie organizacji.

2. Zgodność regulacyjna – AI Act, NIS 2, RODO i przesunięta odpowiedzialność

Drugim filarem odpowiedzialnej transformacji jest pełne osadzenie wdrożeń AI w obowiązujących i wchodzących w życie ramach regulacyjnych. Krajobraz prawny zmienił się w ostatnich latach radykalnie – AI Act, dyrektywa NIS 2, znowelizowana ustawa o krajowym systemie cyberbezpieczeństwa oraz RODO tworzą razem gęstą siatkę obowiązków, w której administracja publiczna funkcjonuje w pierwszym rzędzie podmiotów objętych regulacją. Co istotne, regulator wyraźnie przesunął ciężar odpowiedzialności z poziomu wyłącznie instytucjonalnego na poziom organizacyjny i osobowy.

W praktyce oznacza to wymierne konsekwencje finansowe i osobiste. AI Act przewiduje kary sięgające 35 mln euro lub 7% rocznego obrotu, NIS 2 – do 10 mln euro lub 2% obrotu. Jeszcze istotniejsza jest odpowiedzialność osobowa: kierownictwo podmiotów publicznych może odpowiadać kwotami sięgającymi do 100% miesięcznego wynagrodzenia, a w sektorze prywatnym do 300%. Tymczasem 90% użytkowników AI nie ma świadomości istnienia jakichkolwiek regulacji w tym obszarze, a 70% urzędników wskazuje brak wewnętrznych procedur jako podstawową barierę. Powstaje sytuacja, w której pracownicy każdego dnia podejmują indywidualne decyzje o wykorzystaniu narzędzi AI, których konsekwencje regulacyjne mogą wykraczać daleko poza ich rolę i kompetencje.

Compliance w obszarze AI nie jest przy tym wyłącznie zadaniem prawnym. Wymaga połączenia kompetencji prawnych, technologicznych i operacyjnych – zmapowania, jakie systemy są wykorzystywane, sklasyfikowania ich według poziomów ryzyka definiowanych przez AI Act, ustanowienia mechanizmów nadzoru człowieka nad decyzjami algorytmicznymi oraz zapewnienia minimalizacji danych wymaganej przez RODO przy trenowaniu i użytkowaniu modeli.

- **Rekomendacja:** *Każdy podmiot administracji publicznej powinien w trybie pilnym ustanowić wewnętrzne ramy zarządzania wykorzystaniem AI obejmujące: politykę użycia AI z jednoznacznym katalogiem zastosowań dozwolonych i zabronionych, mapę wykorzystywanych systemów wraz z ich klasyfikacją według poziomów ryzyka AI Act, mechanizmy nadzoru człowieka nad decyzjami administracyjnymi wspieranymi przez AI oraz wyznaczenie ról odpowiedzialnych za zgodność (compliance officer, koordynator AI, zespół ds. cyberbezpieczeństwa). Celem tych ram jest nie tylko ochrona instytucji, ale też realne wsparcie pracowników, którzy w obecnym stanie rzeczy ponoszą indywidualną odpowiedzialność za działania, do których organizacja nie dała im narzędzi ani procedur.*

3. Bezpieczeństwo danych i lokalność stosu technologicznego

Trzeci obszar dotyczy fundamentalnego pytania: gdzie i jak przetwarzane są dane, którymi karmiona jest sztuczna inteligencja w administracji publicznej. Administracja przetwarza najbardziej wrażliwe informacje o obywatelach – od danych medycznych, przez dane finansowe, po informacje dotyczące spraw karnych i administracyjnych. Jednocześnie należy do najbardziej atakowanych sektorów w Polsce: rosnąca liczba ataków ransomware, wycieków danych oraz zupełnie nowych form zagrożeń, takich jak deepfake i AI-asystowana dezinformacja, czyni z bezpieczeństwa cyfrowego jedno z kluczowych wyzwań funkcjonowania państwa. 89% Polaków wskazuje cyberzagrożenia jako poważny problem, a 63,2% wprost domaga się procedur reagowania na cyberataki.

Kluczowym ryzykiem operacyjnym pozostaje masowe wykorzystanie publicznych, darmowych narzędzi AI do pracy z danymi urzędowymi. 90% urzędników korzystających z AI sięga właśnie po takie rozwiązania, często wprowadzając do nich treści obejmujące dane osobowe, dokumenty wewnętrzne lub fragmenty postępowań. Dane te trafiają na serwery zlokalizowane poza Polską i często poza Unią Europejską, gdzie mogą być wykorzystywane do dalszego trenowania modeli, co oznacza utratę kontroli nad informacją. Zjawisko to – określane mianem „shadow AI” – stanowi dziś jedno z najpoważniejszych ryzyk operacyjnych w administracji.

Odpowiedzią musi być budowa zaufanego stosu technologicznego – warstwowego zestawu narzędzi i infrastruktury, w którym dane administracji publicznej są przetwarzane w sposób kontrolowany, lokalny i zgodny z regulacjami. 78% jednostek wskazuje obawę przed uzależnieniem od jednego dostawcy (vendor lock-in) jako istotne ryzyko stra-

tegiczne. Lokalność staje się więc nie tylko kwestią suwerenności, ale też praktycznego bezpieczeństwa operacyjnego – od fizycznej lokalizacji danych, przez certyfikację infrastruktury, po kontrolę nad cyklem życia modelu i danych treningowych.

W praktyce można mówić o wyraźnej hierarchii dostępnych modeli wdrożeniowych – uporządkowanej według poziomu kontroli, jaki organizacja zachowuje nad swoimi danymi. Najwyższy poziom bezpieczeństwa zapewniają rozwiązania działające w pełni lokalnie, instalowane bezpośrednio w infrastrukturze urzędu (on-premise) – w modelu, w którym dane nigdy nie opuszczają granic organizacji, a komunikacja z modelem odbywa się wewnątrz sieci instytucji, bez wysyłania jakichkolwiek treści do serwerów zewnętrznych. Polski rynek oferuje już dziś dojrzałe rozwiązania tego typu, dedykowane sektorowi publicznemu i podmiotom regulowanym, które pozwalają korzystać z generatywnej AI w pełni offline względem internetu publicznego. Wyższym poziomem zaufania od narzędzi globalnych jest również krajowa infrastruktura obliczeniowa, której rozwój Polska konsekwentnie wspiera – między innymi poprzez Fabryki AI w Krakowie i Poznaniu, udział w bałtyckiej Gigafabryce AI oraz Rządowy Klastr Bezpieczeństwa. Najniższy poziom kontroli – i tym samym najwyższy poziom ryzyka – reprezentują darmowe, publicznie dostępne narzędzia globalne, które pozostają dziś najczęstszym wyborem urzędników.

Rekomendacja: Administracja publiczna powinna przyjąć zasadę: dane wrażliwe i urzędowe nie opuszczają zaufanego stosu technologicznego. W praktyce oznacza to wprowadzenie kategorię zakazu wprowadzania danych objętych tajemnicą służbową, danych osobowych oraz fragmen-

tów postępowań do publicznych, niekontrolowanych narzędzi AI; dopasowanie modelu wdrożeniowego do poziomu wrażliwości przetwarzanych danych; oraz – dla najwrażliwszych zastosowań – uczynienie rozwiązań w pełni lokalnych, instalowanych bezpośrednio w infrastrukturze urzędu i działających bez wysyłania danych do jakichkolwiek serwerów zewnętrznych, modelem domyślnym, a nie wyjątkiem. Polskie rozwiązania on-premise dostępne na rynku pozwalają to dziś realnie osiągnąć. Lokalność stosu technologicznego nie jest hasłem politycznym – jest praktycznym warunkiem zgodności z RODO, AI Act i NIS 2 oraz odpornością państwa na ryzyka geopolityczne.

4. Kompetencje, kultura organizacyjna i odpowiedzialne wdrażanie po stronie ludzi

Czwartym obszarem jest najczęściej niedoceniany, a zarazem najbardziej decydujący o powodzeniu transformacji wymiar – wymiar ludzki. Najbardziej zaawansowana technologia nie przyniesie spodziewanych rezultatów, jeśli pracownicy nie będą jej rozumieli, nie będą umieli z niej korzystać w sposób bezpieczny, lub – co dziś realnie się dzieje – będą ukrywać korzystanie z niej przed pracodawcą.

Skala luki kompetencyjnej w polskiej administracji jest znacząca. Mimo że 60% pracowników deklaruje korzystanie z AI w pracy, jedynie 29% przeszło jakiejkolwiek formalne lub nieformalne szkolenie z tego zakresu. Ponad połowa pracowników (55%) nie ujawnia faktu

korzystania z AI, a podobny odsetek (55%) traktuje wynik pracy AI jako własny – co rodzi poważne implikacje regulacyjne i etyczne. Jednocześnie 87% urzędników deklaruje gotowość udziału w dedykowanych szkoleniach z zakresu AI, co stanowi rzadko spotykany kapitał motywacyjny gotowy do uruchomienia.

Drugą stroną problemu jest kultura organizacyjna. W obecnym stanie rzeczy pracownicy często działają w warunkach niepewności – bez jasnych procedur, narzędzi i sygnałów ze strony organizacji, czy korzystanie z AI jest akceptowane, oczekiwane czy zakazane. To prowadzi do sytuacji, w której najambitniejsi i najbardziej innowacyjni urzędnicy ukrywają swoje praktyki, bojąc się sankcji, a organizacja traci możliwość uczenia się na podstawie ich doświadczeń. Skutek jest podwójnie negatywny: ryzyka rosną, bo nikt ich nie monitoruje, a wartość innowacji się rozprasza, bo nie jest dzielona.

Rekomendacja: Administracja publiczna powinna zainwestować w systemowy program rozwoju kompetencji AI, obejmujący trzy poziomy: szkolenia podstawowe dla wszystkich pracowników (zasady bezpiecznego korzystania, granice prawne, ochrona danych), szkolenia praktyczne dla użytkowników operacyjnych (efektywne wykorzystanie narzędzi w konkretnych obszarach pracy) oraz szkolenia strategiczne dla kadry zarządzającej (zarządzanie ryzykiem, compliance, kierowanie zespołami w środowisku z AI). Równoległe organizacje powinny aktywnie zmieniać kulturę z karzącej za niezgłoszone korzystanie

na wspierającą jego mądre, otwarte stosowanie – poprzez kanały zgłaszania wątpliwości, sieci wewnętrznych ambasadorów AI oraz regularne fora wymiany dobrych praktyk między urzędami.

5. Ekosystem partnerstw i wspólna infrastruktura

Piątym i ostatnim obszarem rekomendacji jest świadome budowanie ekosystemu wsparcia, w którym administracja publiczna nie wdraża AI w izolacji. Skala i złożoność transformacji są takie, że żadna pojedyncza instytucja – nawet najbardziej zasobna – nie jest w stanie samodzielnie udźwignąć wszystkich wymiarów zmiany: technologicznego, prawnego, operacyjnego i kompetencyjnego. Próba samodzielnego rozwiązywania każdego problemu od podstaw oznacza marnowanie publicznych środków, powielanie tych samych błędów w różnych urzędach i opóźnianie korzyści, które obywatele mogliby odczuwać już dziś.

Z drugiej strony, korzystanie z doświadczenia rynku oraz wyspecjalizowanych partnerów łączących kompetencje technologiczne, prawne i operacyjne pozwala skrócić czas wdrożeń, ograniczyć ryzyko błędów, korzystać ze sprawdzonych standardów i utrzymywać zgodność z dynamicznie zmieniającymi się regulacjami. Naturalnym katalizatorem takiej współpracy są organizacje branżowe i izby gospodarcze skupiające ekspertów z obszaru cyfryzacji, sztucznej inteligencji i cyberbezpieczeństwa – instytucje, które łączą funkcję merytorycznego zaplecza, dostawcy programów szkoleniowych, autora analiz rynkowych oraz koordynatora dialogu między administracją, biznesem i regulatorem. To one stanowią dziś jedno z niewielu źródeł wiedzy obejmującej jednocześnie wymiar technologiczny, prawny i operacyjny transformacji AI w sektorze publicznym. Równolegle Polska

rozwija własną bazę technologiczną – od krajowej infrastruktury obliczeniowej i polskich modeli językowych, po dostawców lokalnych rozwiązań on-premise, którzy umożliwiają wdrożenia AI bezpośrednio w infrastrukturze urzędu, bez konieczności korzystania z chmury – co tworzy zasób wspólny, możliwy do wykorzystania przez całą administrację publiczną.

Trzecim wymiarem ekosystemu jest wymiana doświadczeń wewnątrz samej administracji. Dziś każdy urząd, każda gmina, każde ministerstwo często rozwiązuje analogiczne problemy w pojedynkę – definiując polityki AI, wybierając narzędzia, projektując szkolenia. Tymczasem w skali kraju rozwiązania jednego urzędu mogą być wzorcem dla sektek innych. Sieć transferu wiedzy, wspólne biblioteki dobrych praktyk, ramowe wzorce dokumentów oraz koordynacja między organami administracji rządowej i samorządowej to zasoby, których wytworzenie nie wymaga ogromnych nakładów, a może dramatycznie przyspieszyć i ujednolicić tempo transformacji.

Rekomendacja: Administracja publiczna powinna aktywnie budować i wykorzystywać cztery warstwy ekosystemu: warstwę partnerów rynkowych – w tym zaufanych dostawców polskich, lokalnych rozwiązań AI on-premise, kancelarii prawnych specjalizujących się w AI Act i NIS 2 oraz integratorów systemów z doświadczeniem w środowiskach regulowanych; warstwę organizacji branżowych i izb gospodarczych – pełniących rolę merytorycznego zaplecza eksperckiego, dostawcy szkoleń, autora raportów i przewodników wdrożeniowych oraz

pośrednika w dialogu z regulatorem; warstwę krajowej infrastruktury technologicznej – Fabryki AI, polskie modele językowe, suwerenna chmura obliczeniowa; oraz warstwę współpracy międzyinstytucjonalnej – wspólne biblioteki dobrych praktyk, ramowe wzorce polityk AI oraz sieci wymiany doświadczeń między urzędami i jednostkami samorządu terytorialnego. Inwestycja w ekosystem nie jest dodatkiem do transformacji – jest jednym z kluczowych warunków jej powodzenia w skali kraju.

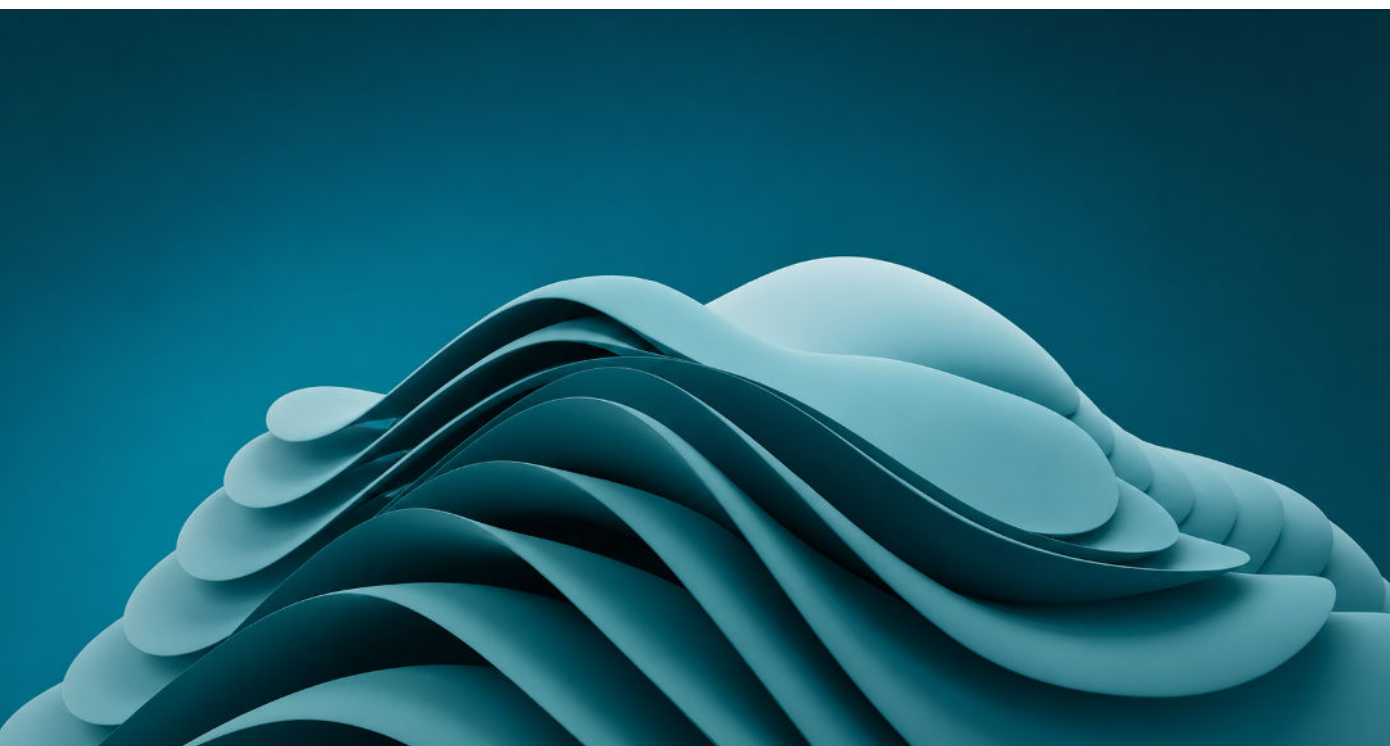
dzie i bezpieczna technologia nie zastąpią wartości płynącej z dzielenia się doświadczeniem w ramach krajowego ekosystemu.

Wspólnym mianownikiem wszystkich rekomendacji jest świadoma rezygnacja z dwóch postaw, które dziś nadal dominują w wielu instytucjach: postawy zakazu, która ignoruje rzeczywiste praktyki pracowników i uniemożliwia kontrolę nad ryzykiem, oraz postawy biernej akceptacji, która delegowanie decyzji o wykorzystaniu AI na pojedynczych urzędników mylnie utożsamia z brakiem konieczności działania. Trzecia droga – droga aktywnego, kontrolowanego i odpowiedzialnego umożliwiania – jest jedyną, która prowadzi do sprawnego państwa, dbającego o najwyższe standardy obsługi obywatela i szanującego jego prywatność.

Podsumowanie – pięć obszarów, jeden kierunek

Pięć przedstawionych obszarów rekomendacji nie jest listą niezależnych zadań, lecz spójnym systemem. Wdrożenia operacyjne nie powiodą się bez ram regulacyjnych. Compliance nie ochroni instytucji bez bezpiecznego stosu technologicznego. Bezpieczny stos nie wystarczy, jeśli pracownicy nie będą umieli z niego korzystać. A kompetentni lu-

Najbliższe lata zdecydują, czy polska administracja publiczna wykorzysta historyczną szansę, jaką daje sztuczna inteligencja, czy też pozostanie na pozycji obserwatora zmian dokonujących się oddolnie i poza jej kontrolą. Wybór należy do decydentów – jednak czas na ten wybór skraca się z każdym kwartałem.





krajowa
izba
gospodarki
cyfrowej

AI W ADMINISTRACJI PUBLICZNEJ

Priorytety Transformacji Cyfrowej w Polsce

Wyzwania i rekomendacje na lata 2026-2029

Perspektywa Obywatela

2026

BĄDŹMY W KONTAKCIE



www.kigc.pl