

Słownik AI

335 haseł z zakresu sztucznej inteligencji, regulacji, suwerenności technologicznej i wdrożeń

Wersja z 21 września 2026 r. • stan prawny: 21.09.2026

Źródło: Słownik pojęć CDF 1.6.1 (Cognitive Deployment Framework) • allclouds.pl sp. z o.o.

Jak korzystać

Hasła ułożone są alfabetycznie. Każde hasło ma etykietę kategorii: Podstawy AI, Regulacje i normy, Suwerenność i architektura, Wdrożenie i zarządzanie lub Metodyka CDF. Wiersz „W regulacji” wskazuje akty prawne i normy, do których hasło się odnosi. Aktualna wersja słownika, wyszukiwarka i kalendarz regulacyjny: www.allclouds.pl/slownik-ai/

Kalendarz regulacyjny

- 17.01.2025** DORA stosowana bezpośrednio w całej UE — DORA
- 2.02.2025** Zakazane praktyki AI i obowiązek kompetencji AI (art. 4) — AI Act
- 2.08.2025** Obowiązki dostawców modeli AI ogólnego przeznaczenia — AI Act
- 7.08.2025** Polska ustawa wdrażająca DORA: KNF kontroluje i nakłada kary — DORA
- 3.04.2026** Nowelizacja KSC wdrażająca NIS2 wchodzi w życie — KSC / NIS2
- 2.08.2026** Przejrzystość (art. 50): informowanie o kontakcie z AI, oznaczanie treści — AI Act
- 3.10.2026** Wpis do wykazu podmiotów kluczowych i ważnych (część z urzędu) — KSC / NIS2
- 28.10.2026** KRiBSI może kontrolować i nakładać kary — Ustawa o systemach SI
- 2.12.2026** Nowe zakazy (treści intymne bez zgody, CSAM); znakowanie treści także w starszych systemach — AI Act
- 3.04.2027** Koniec okresu dostosowawczego KSC: SZBI i pozostałe obowiązki — KSC / NIS2
- 2.12.2027** Systemy wysokiego ryzyka z załącznika III — AI Act
- 3.04.2028** Początek kar KSC i pierwszy audyt nowych podmiotów kluczowych — KSC / NIS2
- 2.08.2028** Wysokie ryzyko w produktach regulowanych (załącznik I) — AI Act

A

A-BCP-01 *AI Service Continuity & DR Plan*

Artefakt CDF (fazy F1.5/F6) mechanizmu CBC – plan ciągłości działania systemu AI wymagany przez art. 21 ust. 2 lit. c NIS2/uKSC. Zasila AIMS Evidence Package oraz pakiet dowodowy uKSC (artefakty K4/K8). Wprowadzony w CDF 1.5.0.

METODYKA CDF

A-F2-16 *Service Fallback Specification*

Artefakt CDF (fazy F2 i F4) kontroli PAG-2 – specyfikacja zastępczej ścieżki kontaktu z pracownikiem dla procesów styku AI z obywatelem. Obejmuje funkcję przełączenia do człowieka, kanał niezależny od systemu AI oraz ścieżkę zgłoszenia zastrzeżeń. Utrzymywana jako sekcja SA-XAI Documentation. Wprowadzony w 1.6.0.

METODYKA CDF

A-GOV-01 *Karta zasadności zastosowania AI*

Artefakt CDF (faza F0) kontroli PAG-1 – udokumentowana odpowiedź na cztery pytania Etapu 0 Przewodnika MC po AI wraz z rozstrzygnięciem GO / NO-GO / REDESIGN i zatwierdzeniem Business Ownera. Wchodzi do Evidence Book jako dowód należytej staranności przy wyborze technologii. Wprowadzony w 1.6.0.

METODYKA CDF

A-GOV-02 *Polityka korzystania z AI w jednostce*

Artefakt CDF (faza F2) kontroli PAG-3 – zasady korzystania z narzędzi AI przez pracowników, katalog danych zakazanych (dane osobowe i identyfikatory, dane szczególnych kategorii, informacje niejawne, wewnętrzne dokumenty urzędowe nieprzeznaczone do upublicznienia), tryb zatwierdzania nowych narzędzi, obowiązki informacyjne. Wprowadzony w 1.6.0.

METODYKA CDF

A-GOV-03 *Rejestr zatwierdzonych narzędzi AI*

Artefakt CDF (fazy F2 i F6) kontroli PAG-3 – wykaz narzędzi AI dopuszczonych do użytku służbowego z podstawą dopuszczenia, zakresem dozwolonych danych, informacją o wykorzystywaniu danych do trenowania modeli na rzecz innych klientów, właścicielem biznesowym i datą przeglądu. Wprowadzony w 1.6.0.

METODYKA CDF

A-HDP-01 *Health Data Protection & EHDS Assessment*

Artefakt CDF (faza F1.5) mechanizmu HDC (kontrola HDC-1) – ocena ochrony danych zdrowotnych oraz zgodności z EHDS i RODO art. 9. Wymagany dla wdrożeń z aktywnym tagiem HEALTH. Wprowadzony w CDF 1.5.0.

METODYKA CDF

A-MDR-01 *Medical Device AI Conformity Dossier*

Artefakt CDF (faza F1.5) mechanizmu HDC (kontrola HDC-2) – dossier zgodności wyrobu medycznego AI (MDR 2017/745 / IVDR 2017/746), generowany gdy system AI kwalifikuje się jako oprogramowanie wyrobu medycznego. Wprowadzony w CDF 1.5.0.

METODYKA CDF

A-MON-01 *Post-Deployment Monitoring Checklist*

Artefakt CDF (A-MON-01) generowany w F5 i utrzymywany w F6, oparty na taksonomii NIST AI 800-4 (marzec 2026). Definiuje 6 kategorii monitorowania zgodnych z NIST: (1) Functionality – czy system działa zgodnie z przeznaczeniem, (2) Operational – czy infrastruktura utrzymuje spójny poziom usługi, (3) Human Factors – jakość interakcji człowiek-AI (wg NIST: rozbieżność intencji i percepcji użytkownika, ciemne wzorce; w CDF dodatkowo automation bias, skill degradation, overtrust), (4) Security – zabezpieczenie przed atakami (OWASP ASI, prompt injection), (5) Compliance – zgodność regulacyjna (AI Act, NIS2, DORA), (6) Large-Scale Impacts – pozytywne efekty społeczne i środowiskowe. Wymagany dla LOA $\geq$ L2. Dla każdej kategorii określa metryki, częstotliwość pomiaru, progi eskalacji i właściciela. Identyfikator artefaktu: A-MON-01.

METODYKA CDF

A-PLD-01 *Product Liability Evidence Dossier*

Artefakt CDF (fazy F1.5/F2/F6) mechanizmu LER – dossier gotowości dowodowej zawierające mapę ról (producent / importer / dystrybutor / znaczący modyfikujący), rejestr znaczących modyfikacji (LER-2), indeks dowodów podlegających ujawnieniu (art. 9 dyrektywy PLD 2024/2853) z klauzulą ochrony tajemnicy przedsiębiorstwa oraz dokumentację adekwatności

bezpieczeństwa produktu (art. 7 ust. 2 lit. f). Zasila Evidence Book (EB-4). Wprowadzony w CDF 1.5.0.

METODYKA CDF

A-SSI-01 *Mapa organów nadzoru AI*

Identyfikacja właściwego organu nadzoru (Komisja, KNF, regulator sektorowy) dla danego wdrożenia AI na terytorium RP. Obejmuje mapowanie art. 5–47 ustawy o systemach AI (organ nadzoru i rozdział 2) na strukturę organizacyjną klienta oraz wskazanie obowiązków informacyjnych i rejestrowych wobec właściwych organów. Obowiązkowy artefakt fazy F0/F1 w wariantcie GOV.PL.

METODYKA CDF

A-SSI-02 *Procedura zgłaszania poważnych incydentów*

Workflow zgłaszania poważnych incydentów do Komisji zgodnie z AI Act art. 73; odrębnie obejmuje wczesne ostrzeżenia i zgłoszenia NIS2/uKSC. Ustawa o systemach AI nie ustanawia odrębnego polskiego kanału zgłoszenia 24h. Definiuje kryteria kwalifikacji, szablon zgłoszenia, eskalację wewnętrzną i zewnętrzną (Komisja, CSIRT GOV/MON/NASK) oraz rejestr statusów z dashboardem monitoringowym. Faza F3/F5.

METODYKA CDF

A-SSI-03 *Checklist gotowości do kontroli*

Lista kontrolna gotowości organizacji na kontrolę Komisji (rozdział 3, art. 48–57), z zawiadomieniem 7 dni przed kontrolą. Obejmuje kompletność dokumentacji technicznej w języku polskim, dostęp do dokumentów, systemów i chmury obliczeniowej (art. 52), wyznaczenie osoby kontaktowej, dostępność logów systemowych i Agent Registry. Uwzględnia rygor natychmiastowej wykonalności decyzji (art. 108). Faza F4.

METODYKA CDF

A-SSI-04 *Szablon opinii indywidualnej*

Wniosek do Komisji o wydanie opinii indywidualnej dotyczącej stosowania AI Act do konkretnego wdrożenia AI (art. 8–18). Opłata 150 zł wynika z art. 11; termin wynosi 30 dni albo 60 dni w sprawie szczególnie skomplikowanej (art. 12), a brak opinii oznacza opinię zgodną ze stanowiskiem wnioskodawcy. Opinia jest wiążąca w zakresie sprawy (art. 13) i po anonimizacji publikowana w BIP (art. 17). Faza F1/F6.

METODYKA CDF

A-SSI-05 *Procedura obsługi kontroli*

Pełny workflow obsługi kontroli Komisji (art. 48–57): zawiadomienie 7 dni → udostępnienie dokumentacji i infrastruktury → protokół kontroli → prawo do zastrzeżeń → zalecenia pokontrolne → plan realizacji zaleceń → raport wykonania. Zalecenia wyznaczają co do zasady terminy nie krótsze niż 30 dni (art. 57); w razie bezpośredniego ryzyka Komisja może nakazać zaprzestanie stosowania lub wycofanie systemu (art. 63) i nadać decyzji rygor natychmiastowej wykonalności (art. 108). Faza F5.

METODYKA CDF

A-SSI-06 *Szablon zgłoszenia do piaskownicy*

Wniosek o dopuszczenie do piaskownicy regulacyjnej AI prowadzonej przez Komisję (rozdział 7, art. 91–103). Udział trwa 6–12 miesięcy i następuje w drodze konkursu; dla MŚP jest nieodpłatny. Uczestnik składa raport roczny (art. 98), a w terminie 90 dni od potwierdzenia uczestnictwa może wnioskować o opinię indywidualną (art. 101). Faza F0/F2.

METODYKA CDF

A-SSI-07 *Matryca kar i sankcji SSI*

Matryca ryzyka kar administracyjnych i sankcji karnych według ustawy o systemach AI (art. 104–114). Kwoty kar wynikają z art. 99 AI Act, a ich równowartość w złotych oblicza się według kursu NBP z 28 stycznia (art. 104 ust. 4). Obejmuje układ 20–70% / 30–90% (art. 70 i 84), obniżkę 10–50% (art. 107) oraz przepisy karne art. 112–113. Faza F1.

METODYKA CDF

Acceptance Criteria Checklist (F0-AC-01) *Acceptance Criteria Checklist*

Artefakt Fazy F0 (Discovery) definiujący kryteria formalne, które muszą być spełnione, aby Klient zaakceptował wynik wdrożenia. Zawiera minimum 6 kryteriów rekomendowanych przez CDF (Cognitive SLA, Coverage Score, Agent Coordination, AIMS Evidence Package, Risk Register, Knowledge Transfer) oraz kryteria własne Klienta. Stanowi załącznik do umowy wdrożeniowej. Identyfikator: F0-AC-01.

METODYKA CDF

Accountability Matrix *Matryca odpowiedzialności per LOA*

Tabela przypisująca odpowiedzialność za decyzje systemu AI do konkretnych ról organizacyjnych w zależności od poziomu autonomii agenta (LOA L0–L4). Dla L0–L1 pełna odpowiedzialność leży po stronie operatora; dla L2 odpowiedzialność jest dzielona; dla L3–L4 odpowiedzialność przesuwa się na właściciela systemu z obowiązkowym audytem. Patrz: sekcja K6/K10 w Metodyce.

METODYKA CDF

ACE (Adaptive Configuration Engine) *Adaptive Configuration Engine*

Dynamiczny mechanizm konfiguracyjny wbudowany w CDF, który aktywuje odpowiedni podzbiór faz, kontroli, artefaktów i załączników na podstawie profilu organizacji i wybranych wymogów zgodności. Działa w trybie Self-Service (Kwestionariusz Konfiguracyjny) lub Guided (F0). Wprowadzony w CDF 1.4.4, rozdział 11.

METODYKA CDF

ACE Configuration Profile *Profil konfiguracyjny ACE*

Dokument wynikowy ACE określający aktywne tagi, wybrany profil referencyjny, aktywne tiers mechanizmów CDF (Cognitive SLA, HCG, F6) oraz spersonalizowaną Matrycę Zgodności. Generowany w F0, wersjonowany jako ACE-[PROFIL]-v[N].[M].

METODYKA CDF

ACE Profile Matrix *Matryca profili ACE*

Załącznik E-ACE: matryca 77 wymogów zgodności × 11 profili referencyjnych z przypisaniem WYMAGANE / REKOMENDOWANE / DOSTĘPNE per wymóg i profil.

METODYKA CDF

ACE Profile Version Record *Rekord wersji profilu ACE*

Formalny, niezmienny (immutable) rekord każdej wersji profilu konfiguracyjnego organizacji. Zawiera wersję, datę, autora, uzasadnienie zmiany i flagę REDUCTION (jeśli Scope Reduction). Schemat: ACE- [PROFIL]-v[N].[M].

METODYKA CDF

ACE Reconfiguration Log *Rejestr rekonfiguracji ACE*

Artefakt CDF (F0–F6) dokumentujący każdą zmianę profilu konfiguracyjnego organizacji: wersja przed/po, data, uzasadnienie, autor, flaga REDUCTION.

METODYKA CDF

Agent (AI) *Agent AI*

Program, który na podstawie modelu językowego samodzielnie planuje i wykonuje ciąg działań — pobiera dane, wywołuje narzędzia, podejmuje decyzje — w granicach uprawnień nadanych przez organizację. Od asystenta różni go to, że działa, a nie tylko odpowiada

PODSTAWY AI

AGENT (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany, gdy wdrożenie obejmuje agenty AI (pojedyncze lub wieloagentowe). Włącza sekcje Agent Governance, tożsamości agentów, bezpieczeństwa i cyklu życia agentów.

METODYKA CDF

Agent Coordination

Metryka mierząca skuteczność współpracy agentów AI w środowisku multi-agent. Definiowana jako procent zadań multi-agent ukończonych bez eskalacji do człowieka. Target bazowy CDF: ≥85%. Metryka należy do Cognitive SLA Tier 2 (6 metryk). Mierzona w Fazie F4 (Cognitive Sprint) i monitorowana w Fazie F6 (CogOps). Patrz: Tabela 11 w Metodyce.

METODYKA CDF

Agent Identity Schema *Agent Identity Schema (CDF, na bazie NIST NCCoE)*

Schemat tożsamości agenta AI. Konstrukcja własna CDF rozwijająca filary identyfikacji i uwierzytelniania z concept paper NIST NCCoE „Accelerating the Adoption of Software and AI Agent Identity and Authorization” (5.02.2026). Nie jest cytatem z publikacji NIST. Kontrola T-011, faza F2, artefakt A-F2-05 Agent Registry.

METODYKA CDF

Agent Kill-Switch *Wyłącznik awaryjny agenta*

Mechanizm natychmiastowego zatrzymania pojedynczego agenta AI albo całego roju agentów — sprzętowy lub programowy, niezależny od samego agenta. Po jego uruchomieniu zależne systemy dostają powiadomienie, a proces przechodzi na weryfikację przez człowieka

PODSTAWY AI

Agent Lifecycle Management (ALM) *Zarządzanie cyklem życia agenta*

Sześciostopniowy model zarządzania agentem w CDF: Design → Build → Test → Deploy → Monitor → Optimize/Retire. Obejmuje także Agent Retirement Protocol zapewniający bezpieczne wycofanie agenta, sprawdzenie zależności oraz transfer wiedzy operacyjnej i biznesowej. Realizowane w Fazie 6 CDF. Zob. ALM (Agent Lifecycle Management).

METODYKA CDF

Agent Readiness Audit *Audyt gotowości na agenty AI*

Ocena dojrzałości organizacji do wdrożenia autonomicznych agentów AI, obejmująca dane, procesy, kompetencje, governance, gotowość psychologiczną i zdolność do bezpiecznego nadzorowania agentów w środowisku pracy. Realizowany w Fazie 0 CDF.

METODYKA CDF

Agent Registry *Rejestr agentów AI*

Obowiązkowy od Fazy 2 centralny rejestr wszystkich agentów systemu, zawierający ich rolę, uprawnienia, poziom autonomii, właściciela, zależności, budżety tokenowe, pole Delegated By, Kill-Switch Authority i historię audytów. Fundament Agent Governance Model oraz audytowalności systemu wieloagentowego.

METODYKA CDF

Agent Retirement Protocol

Formalna procedura Fazy 6 CDF definiująca bezpieczne wycofanie agenta AI z eksploatacji. Obejmuje: (1) archiwizację pamięci i historii decyzji agenta, (2) cofnięcie uprawnień w Agent Registry, (3) sprawdzenie zależności – czy inne agenty lub procesy polegają na wycofywanym agencie, (4) transfer wiedzy operacyjnej do następcy lub zespołu ludzkiego, (5) weryfikację braku aktywnych subskrybentów MCP, (6) raport końcowy z powodami i analizą wpływu.

METODYKA CDF

Agent Threat Model

Spersonalizowany model zagrożeń agentowych oparty na CoSAI, OWASP ASI i CSA ATF. Filtrowany przez profil ACE. Patrz: Załącznik B.

METODYKA CDF

AI *Sztuczna inteligencja (Artificial Intelligence)*

Sztuczna inteligencja — systemy komputerowe zdolne do wykonywania zadań, które dotąd wymagały ludzkiej inteligencji: rozumowania, uczenia się, klasyfikacji, generowania treści i podejmowania decyzji. W prawie UE definicję systemu AI zawiera art. 3 pkt 1 AI Act

PODSTAWY AI · W regulacji: AI Act

AI Act *Akt w sprawie sztucznej inteligencji — rozporządzenie (UE) 2024/1689*

Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 ustanawiające zharmonizowane przepisy dotyczące sztucznej inteligencji, znowelizowane rozporządzeniem (UE) 2026/1744 (Digital Omnibus on AI). Klasyfikuje systemy AI według poziomu ryzyka i nakłada obowiązki na dostawców oraz podmioty wdrażające: zarządzanie ryzykiem, rejestrowanie zdarzeń, przejrzystość, nadzór człowieka i dokumentację techniczną. Stosowanie etapowe: zakazane praktyki i kompetencje w zakresie AI od 02.02.2025; modele AI ogólnego przeznaczenia, nadzór i kary od 02.08.2025; obowiązki przejrzystości (art. 50) od 02.08.2026; nowe zakazy z art. 5 od 02.12.2026; samodzielne systemy wysokiego ryzyka z załącznika III od 02.12.2027; systemy wbudowane w produkty z załącznika I od 02.08.2028

REGULACJE I NORMY · W regulacji: rozp. 2026/1744

AI Act Readiness Score (ARS) *AI Act Readiness Score*

Ilościowy wskaźnik (0–100%) gotowości organizacji do zgodności z EU AI Act. Formuła: $ARS = \sum (waga \times wynik)$ dla 5 wymiarów: Risk Classification, Technical Documentation, Human Oversight, Transparency, FRIA Readiness. Progi decyzyjne: $\geq 90\%$ – pełna zgodność (full compliance), 70–89% – zgodność warunkowa (conditional), $< 70\%$ – blokada wdrożenia (blocked). Generowany w F0 i aktualizowany w F1.5. Sformalizowany w CDF 1.4.4 (Z-003).

METODYKA CDF

AI Champions Program *Program liderów AI (AI Champions)*

Program budowania wewnętrznych liderów transformacji AI: przeszkolonych pracowników, którzy w swoich zespołach wspierają adopcję narzędzi, eskalują problemy i promują dobre praktyki korzystania z systemów AI. Skuteczniejszy niż szkolenia masowe, bo zmiana idzie od osób, którym zespół ufa

WDROŻENIE I ZARZĄDZANIE

AI Concerns Feedback Channel *Kanał zgłaszania obaw dotyczących AI*

Formalny mechanizm, dzięki któremu osoby, których dotyczy działanie systemu AI — pracownicy, klienci, obywatele — mogą zgłaszać obawy, skargi i wnioski dotyczące decyzji podejmowanych z udziałem AI. Wymagany przez ISO/IEC 42001 (załącznik A, kontrola A.3.3) i wspierający nadzór człowieka z art. 14 AI Act

WDROŻENIE I ZARZĄDZANIE · W regulacji: AI Act · ISO/IEC 42001

AI Fatigue Index *Indeks zmęczenia AI*

Wskaźnik mierzący zmęczenie organizacji tempem wdrożeń AI: przeciążenie zmianą, nadmiar równoległych inicjatyw, spadek zaangażowania użytkowników. Pozwala dozować skalę transformacji, zanim entuzjazm zamieni się w bierny opór

WDROŻENIE I ZARZĄDZANIE

AI Model & Vendor Risk Assessment

Kompleksowy proces oceny ryzyka modeli AI i ich dostawców realizowany w Fazie 1.5 CDF. Obejmuje 5 kryteriów: (1) jurysdykcja dostawcy i lokalizacja przetwarzania, (2) dane treningowe i prawa własności, (3) compliance disclosure i certyfikacje, (4) ryzyko vendor lock-in i dostępność alternatyw (VDP-1), (5) historia regulacyjna i incydenty bezpieczeństwa. Wynik dokumentowany w Tabeli nr 8 Metodyki.

METODYKA CDF

AI Vendor Diversification Strategy

Strategia minimalizacji ryzyka uzależnienia od pojedynczego dostawcy AI (VDP-1).

METODYKA CDF

AIMS *AI Management System — system zarządzania sztuczną inteligencją (ISO/IEC 42001)*

System zarządzania sztuczną inteligencją zdefiniowany w normie ISO/IEC 42001:2023 — zbiór polityk, ról, procesów i kontroli, dzięki którym organizacja rozwija i wykorzystuje AI w sposób nadzorowany, udokumentowany i możliwy do audytu. Odpowiednik systemu zarządzania bezpieczeństwem informacji (ISO/IEC 27001) dla obszaru AI; podlega certyfikacji

REGULACJE I NORMY · W regulacji: ISO/IEC 42001

AIMS Evidence Package

Zbiór dowodów zgodności z ISO 42001 generowany i utrzymywany w ramach CDF. Obejmuje: SoA, ASIA Report, AI Policy, Agent Registry, Compliance Readiness Report. Tworzony od F1.5, enriched w F5, utrzymywany w F6.

METODYKA CDF

AIMS Lifecycle *Cykl życia artefaktów AIMS*

Tabela definiująca cykl życia 12 kluczowych artefaktów pakietu AIMS Evidence Package w fazach CDF F0–F6: moment utworzenia (Created), aktualizacji (Updated), przeglądu (Reviewed) i właściciela (Owner). Źródło: Metodyka, sekcja AIMS Evidence Package. Wprowadzone w v1.4.4.

METODYKA CDF

AIP *Agent Interaction Protocols*

Zestaw protokołów zarządzających współpracą agentów AI w CDF. Obejmuje trzy wzorce: Hierarchical (Supervisor-Worker, domyślny dla procesów krytycznych L0-L1), Collaborative (współdzielona przestrzeń, L2-L3) i Pipeline (sekwencyjne przekazywanie wyników). Każda interakcja maszyna-maszyna jest logowana w Immutable Audit Trail.

METODYKA CDF

Air-gap *Air gap — izolacja sieciowa*

Fizyczne odcięcie systemu od sieci zewnętrznych, w tym internetu. Najwyższy poziom izolacji dla suwerennych systemów AI, typowy dla obronności, informacji niejawnych i infrastruktury krytycznej. Wymaga osobnych procedur aktualizacji modeli i przenoszenia danych, bo nic nie wchodzi do środowiska ani z niego nie wychodzi bez kontroli

SUWERENNOŚĆ I ARCHITEKTURA

Air-Gapped Defense *Izolowana architektura obronna (air-gapped)*

Wzorzec wdrożenia AI zakładający pełną fizyczną izolację środowiska: brak połączeń z internetem i brak operacyjnych zależności od usług zewnętrznych — modele, dane, logi i orkiestracja działają wyłącznie wewnątrz. Stosowany w obronności, przy informacjach niejawnych i w wybranych systemach infrastruktury krytycznej państwa

SUWERENNOŚĆ I ARCHITEKTURA

ALM (Agent Lifecycle Management) *Agent Lifecycle Management*

Zarządzanie pełnym cyklem życia agenta AI: od rejestracji w Agent Registry (F2), przez monitoring operacyjny (F4-F6), do wycofania (Agent Retirement Protocol). Zob. Agent Lifecycle Management (ALM).

METODYKA CDF

API *Application Programming Interface — interfejs programistyczny*

Interfejs programistyczny — zestaw reguł i metod, dzięki którym systemy lub komponenty oprogramowania komunikują się ze sobą. W systemach AI każde API (do modelu, narzędzi, danych) jest elementem łańcucha dostaw i podlega ocenie bezpieczeństwa integracji

PODSTAWY AI

ARS (AI Act Readiness Score) *AI Act Readiness Score*

Skrót od AI Act Readiness Score. Zob. AI Act Readiness Score (ARS).

METODYKA CDF

ASIA *AI System Impact Assessment*

Ocena wpływu systemu AI realizowana przed wdrożeniem, obejmująca analizę ryzyk dla praw podstawowych, bezpieczeństwa, środowiska i społeczeństwa. W CDF inicjowana w Fazie F1.5 (dla profili z tagiem REGULATED) lub aktywowana przez profil ACE (REKOMENDOWANE dla pozostałych). Aktualizowana w kolejnych fazach (F2, F5, F6). Mapuje się na obowiązki art. 9 i art. 27 EU AI Act oraz na wymagania FRIA. Dla wdrożeń CSE-GOV i CSE-COM z danymi regulowanymi raport zawiera rozszerzony moduł suwerenności.

METODYKA CDF

ATF (Agentic Trust Framework) *CSA Agentic Trust Framework*

Framework dojrzałości agentowej opublikowany jako otwarta specyfikacja na blogu Cloud Security Alliance (2.02.2026; autor: J. Woodruff, MassiveScale.AI – nie jest formalną publikacją CSA). 4 poziomy dojrzałości: Level 1 (Intern – obserwacja, zatwierdza każdą akcję), Level 2 (Junior – rekomenduje i wykonuje po zatwierdzeniu), Level 3 (Senior – działa samodzielnie z powiadomieniem ex-post), Level 4 (Principal – pełna autonomia w granicach mandatu). Mapowanie na CDF LOA: L0-L1 (Intern), L2 (Junior), L3 (Senior), L4 (Principal). Patrz: Tabela 57 (ATF ↔ LOA Cross-Reference) w Załączniku B Metodyki.

METODYKA CDF

Autonomous Swarm *Autonomiczny rój agentów*

Sieć wielu współpracujących agentów AI, które działają bez stałego nadzoru człowieka w zdefiniowanych granicach autonomii. Stosowana w procesach zaplecza o niskiej krytyczności, zawsze z limitem kosztów i wyłącznikiem awaryjnym

PODSTAWY AI

B

Backup 3-2-1 *Strategia kopii zapasowych 3-2-1*

Zasada tworzenia kopii zapasowych: co najmniej trzy kopie danych, na dwóch różnych nośnikach, z czego jedna poza główną lokalizacją. W systemach AI obejmuje nie tylko dane, lecz także wagi modeli, indeksy wektorowe, konfiguracje i logi — bez nich odtworzenie systemu nie jest możliwe

SUWERENNOŚĆ I ARCHITEKTURA

Badanie NASK 2026 *AI w e-administracji publicznej — badanie NASK (2026)*

Raport NASK – PIB z 2026 r. o wykorzystaniu AI w administracji publicznej z perspektywy urzędników i instytucji. Autorzy określają badanie jako eksploracyjne i pilotażowe, a próbę jako niereprezentatywną, dlatego wyniki należy traktować jako materiał porównawczy, a nie parametry dla całej administracji ani podstawę zobowiązań umownych

REGULACJE I NORMY

Biblioteka wzorców wdrożeniowych *Biblioteka wzorców wdrożeniowych AI*

Skatalogowany zbiór sprawdzonych wzorców wdrożeń AI wielokrotnego użytku — architektur, konfiguracji, promptów, procedur odbioru — pozwalający szybko i bezpiecznie powtarzać udane wdrożenia w kolejnych działach lub jednostkach organizacji

WDROŻENIE I ZARZĄDZANIE

Biuro Komisji *Obsługa Komisji Rozwoju i Bezpieczeństwa Sztucznej Inteligencji*

Ustawa o systemach sztucznej inteligencji nie tworzy odrębnego biura ani państwowej osoby prawnej: obsługę administracyjną Komisji i jej Przewodniczącego zapewnia urząd obsługujący

ministra właściwego do spraw informatyzacji. Wydatki Komisji są finansowane z budżetu państwa w granicach limitów określonych w ustawie, a wpływy z administracyjnych kar pieniężnych stanowią dochód budżetu państwa

REGULACJE I NORMY

Bounded Autonomy *Ograniczona autonomia agenta*

Zasada projektowa, zgodnie z którą każdy agent AI działa w ściśle określonych granicach uprawnień, budżetu (tokenów, kosztów) i zakresu decyzji. Granice przypisuje się na podstawie skali autonomii (L0-L4) i stale monitoruje

PODSTAWY AI

Business Continuity *Ciągłość działania (business continuity)*

Zdolność organizacji do utrzymania kluczowych funkcji w trakcie zakłócenia i po nim — awarii, ataku, utraty dostawcy. Dla systemów AI oznacza m.in. plan działania bez dostępu do modelu lub API zewnętrznego, możliwość wyłączenia agentów i przywrócenia pracy ręcznej

SUWERENNOŚĆ I ARCHITEKTURA

C

C-Suite *C-suite — najwyższe kierownictwo*

Zbiórce określenie najwyższego kierownictwa organizacji — zarządu i dyrektorów z tytułami „chief” (CEO, CFO, CTO, CIO, CISO). W transformacji AI to poziom, na którym zapadają decyzje o priorytetach, budżecie i akceptowalnym ryzyku

WDROŻENIE I ZARZĄDZANIE

CapEx *CapEx — wydatki inwestycyjne*

Wydatki kapitałowe na zakup lub budowę aktywów trwałych — w projektach AI m.in. serwerów, akceleratorów GPU, pamięci masowej, chłodzenia i okablowania. Zestawiane z wydatkami operacyjnymi (OpEx) w modelu całkowitego kosztu posiadania na 3-5 lat

WDROŻENIE I ZARZĄDZANIE

CBC-1 *Continuity & DR Plan*

Kontrola CDF mechanizmu CBC: plan ciągłości działania i przywracania (BCP/DRP) dla usług AI – RTO/RPO per system, kopie zapasowe 3-2-1, tryby kontrolowanej degradacji, fallback na model suwerenny (spięcie z VDP-2). Obszary wiedzy KD5/KD11. Fazy F1/F1.5/F6. Wprowadzona w CDF 1.5.0.

METODYKA CDF

CBC-2 *Crisis Management & Comms*

Kontrola CDF mechanizmu CBC: procedury zarządzania kryzysowego i komunikacji, eskalacja oraz powiązanie z obsługą incydentów (24h/72h wg NIS2/uKSC). Obszary wiedzy KD5/KD2. Fazy F1.5/F6. Wprowadzona w CDF 1.5.0.

METODYKA CDF

CBC-3 *Continuity Testing*

Kontrola CDF mechanizmu CBC: cykliczne testy odtworzeniowe oraz utrzymanie aktualności planu ciągłości działania. Obszary wiedzy KD5/KD12. Faza F6. Wprowadzona w CDF 1.5.0.

METODYKA CDF

CBP *Cyfrowy Bliźniak Poznawczy (Cognitive Worker Twin)*

Cyfrowy model wiedzy i sposobu pracy konkretnego pracownika lub zespołu: wiedza jawna i ukryta, interpretacje, preferencje decyzyjne, wzorce zachowań. Pozwala asystentowi AI odpowiadać tak, jak odpowiedziałby ekspert organizacji, a organizacji — zachować wiedzę, gdy ekspert odchodzi

PODSTAWY AI

CBP Fidelity Test

Test wierności poznawczej realizowany w F5, weryfikujący czy CBP w produkcji zachowuje się zgodnie ze specyfikacją z fazy pilotażowej.

METODYKA CDF

CBP Monitoring Dashboard

Artefakt F5/F6 wizualizujący metryki CBP w czasie rzeczywistym. Oparty na taksonomii NIST AI 800-4.

METODYKA CDF

CBP Readiness Assessment

Artefakt F1 oceniający gotowość procesów biznesowych do kognitywizacji.

METODYKA CDF

CBP Sync Protocol *Protokół Synchronizacji CBP*

Mechanizm ciągłej aktualizacji Cyfrowego Bliźniaka Poznawczego z nowych danych pracownika. Artefakt Fazy F6 (Operate) zapewniający, że CBP pozostaje aktualny w miarę ewolucji wiedzy i zachowań pracownika.

METODYKA CDF

CDF *CDF — Cognitive Deployment Framework*

Metodyka allclouds.pl wdrażania systemów kognitywnych w przedsiębiorstwach i administracji, niezależna od konkretnych technologii i infrastruktury. W edycji Client Success dostosowana do sektora publicznego i regulowanego. Obejmuje pełny cykl: od rozpoznania i zgodności regulacyjnej, przez wdrożenie i skalowanie, po nadzór i ciągłe doskonalenie

WDROŻENIE I ZARZĄDZANIE

CDF Academy *Akademia CDF*

Program szkoleniowy CDF (Faza 3) dostarczający 40-godzinne certyfikujące szkolenie z modułem projektowania interakcji minimalizujących obciążenie poznawcze operatorów. Fundament budowy wewnętrznych kompetencji AI i programu AI Champions.

METODYKA CDF

CDF North Star *Dokument North Star CDF*

Artefakt Fazy 1 łączący cele strategiczne, KPI i kamienie milowe programu AI.

METODYKA CDF

CDF Platform

Aplikacja webowa wspierająca wdrożenie metodyki CDF: interaktywny konfigurator profilu organizacji, który na podstawie sektora, klasyfikacji danych i ekspozycji regulacyjnej wskazuje wymagane kontrole, artefakty i matrycę zgodności. Jedna z warstw Suwerennego Ekosystemu AI allclouds.pl

WDROŻENIE I ZARZĄDZANIE

CDF Risk Engine *Silnik zarządzania ryzykiem*

Rozwiązanie technologiczne wdrażane jako fundament bezpieczeństwa systemu, realizujące ciągłe procedury zarządzania ryzykiem zgodne z ISO/IEC 23894 i wymaganiami governance AI.

METODYKA CDF

Center of Excellence *Centrum kompetencji AI (Center of Excellence)*

Wewnętrzna jednostka organizacyjna odpowiedzialna za standardy, narzędzia, zarządzanie wiedzą i skalowanie rozwiązań AI w całej organizacji. Skupia kompetencje, które w pojedynczych działach byłyby rozproszone lub dublowane

WDROŻENIE I ZARZĄDZANIE

CETS 225 (Konwencja Rady Europy o AI) *Konwencja ramowa Rady Europy o sztucznej inteligencji (CETS 225)*

Konwencja ramowa Rady Europy o sztucznej inteligencji, prawach człowieka, demokracji i praworządności — pierwszy wiążący traktat międzynarodowy o AI, otwarty do podpisu 5 września 2024 r. w Wilnie. Unia Europejska złożyła dokument zatwierdzenia 15 maja 2026 r. i jest dotąd jedyną stroną, która to zrobiła. Konwencja wejdzie w życie po pięciu ratyfikacjach, w tym trzech państw członkowskich Rady Europy; ratyfikacja Unii nie zalicza się do tej trójki

REGULACJE I NORMY

CFL *Cognitive Fidelity Levels*

Pięciostopniowa skala dojrzałości CBP, niezależna od skali LOA. Mierzy wierność odwzorowania poznawczego: CFL-1 Profil Deklaratywny, CFL-2 Profil Obserwacyjny, CFL-3 Profil Interpretacyjny (CFS $\geq 60\%$), CFL-4 Profil Predykcyjny (CFS $\geq 80\%$), CFL-5 Profil Autonomiczny (CFS $\geq 90\%$).

METODYKA CDF

CFS *Cognitive Fidelity Score*

Miara jakości odwzorowania poznawczego w skali 0-100%. Określa, w jakim stopniu Cyfrowy Bliźniak Poznawczy (CBP) wiernie reprezentuje sposób myślenia, interpretacji i podejmowania decyzji przez pracownika. Progi: CFL-3 $\geq 60\%$, CFL-4 $\geq 80\%$, CFL-5 $\geq 90\%$.

METODYKA CDF

Chain of Custody *Łańcuch odpowiedzialności (chain of custody) w systemie AI*

Pełna, udokumentowana ścieżka od danych wejściowych do decyzji wyjściowej systemu AI: każdy etap przekształcenia, wzbogacenia i wykorzystania danych oraz to, kto lub co je zatwierdził. Jest szersza niż pochodzenie danych (data lineage): obejmuje także decyzje ludzi (kto, na jakiej podstawie) i agentów (który agent, z jakim poleceniem, w jakiej roli)

SUWERENNOŚĆ I ARCHITEKTURA

Champion (AI Champion) *Lider AI (AI champion)*

Pracownik organizacji przeszkolony jako ekspert i promotor wdrożeń AI w swoim zespole. Pomaga kolegom w praktycznym użyciu narzędzi i jest pomostem między użytkownikami, biznesem a zespołem transformacyjnym

WDROŻENIE I ZARZĄDZANIE

Change Management *Zarządzanie zmianą*

Systematyczne prowadzenie organizacji przez transformację — od planowania, komunikacji i szkoleń po monitorowanie adopcji i pracę z oporem. We wdrożeniach AI równie ważne jak technologia: to ludzie decydują, czy narzędzie będzie używane

WDROŻENIE I ZARZĄDZANIE

CLASSIFIED (tag ACE) *Tag ACE dla wdrożeń niejawnych*

Tag ACE stosowany w profilach obejmujących systemy przetwarzające informacje niejawne. Aktywuje wymagania związane z akredytacją bezpieczeństwa (np. ABW), fizyczną separacją infrastruktury, kryteriami personalnymi oraz zgodnością z przepisami o ochronie informacji niejawnych. Dotyczy profili DEFENSE i SOV (Sovereign) na najwyższym poziomie restrykcyjności.

METODYKA CDF

Cognitive Circuit Breaker *Kognitywny wyłącznik automatyczny*

Mechanizm automatycznego odcięcia agenta od systemów zewnętrznych po przekroczeniu 3 kolejnych obłanych testów weryfikacji jakości kognitywnej. Zapobiega propagacji błędnych decyzji i stanowi zabezpieczenie dla środowisk wieloagentowych.

METODYKA CDF

Cognitive Degradation Monitoring *Monitoring degradacji kognitywnej*

Trójstopniowa skala monitorowania jakości rozumowania agentów AI: Normal (działanie w normie), Degraded (obniżona jakość – wymagana interwencja), Critical (krytyczne odchylenie – procedury naprawcze). Wdrożona w Fazie 2, generuje alerty i uruchamia procedury naprawcze, a wyniki prezentowane są na dashboardzie CogOps. Uwaga: skala opisuje stany operacyjne agenta, niezależna od skali pomiaru metryk (Cognitive SLA Tiers).

METODYKA CDF

Cognitive Observability *Obserwowalność kognitywna*

Zdolność do szczegółowego monitorowania i logowania ścieżek rozumowania AI, działań agentów, zależności między nimi i wykorzystanych źródeł wiedzy — tak, by proces decyzyjny dało się odtworzyć po fakcie. Wspiera audyt, zgodność, diagnostykę błędów i poprawę jakości decyzji

PODSTAWY AI

Cognitive Quality Reports *Raporty jakości kognitywnej*

Comiesięczne raporty generowane w Fazie 4 i kolejnych, prezentujące wyniki metryk Cognitive SLA, trendy jakości rozumowania, incydenty kognitywne i rekomendacje optymalizacyjne.

METODYKA CDF

Cognitive SLA *Umowa o poziomie usług kognitywnych*

Zestaw metryk jakości rozumowania systemu AI, wykraczający poza klasyczne wskaźniki IT (dostępność, czas odpowiedzi). Mierzy m.in. trafność odpowiedzi, poziom halucynacji, świeżość wiedzy, spójność działania agentów i czas reakcji na wykryty błąd — i wiąże je z progami, po których uruchamia się eskalacja

PODSTAWY AI

Cognitive SLA Tier *Poziom Cognitive SLA*

Trójpoziomowy podział mechanizmu Cognitive SLA: Tier 1 (CORE – 4 metryki bazowe: System Availability, Reasoning Accuracy Rate, Hallucination Rate, Mitigation Response Time), Tier 2 (AGENT + PRODUCTION – +2 metryki koordynacji: Agent Coordination, Confidence Calibration; łącznie 6 metryk), Tier 3 (SCALE lub SOVEREIGN – +Knowledge Freshness Index; łącznie 7 metryk + dashboardy CogOps). Każdy tier jest w pełni operacyjny na swoim poziomie.

METODYKA CDF

Cognitive Sprint *Sprint kognitywny (cognitive sprint)*

Iteracyjny cykl wdrożeniowy w metodyce CDF, analogiczny do sprintu w metodykach zwinnych, ale nastawiony na szybkie i kontrolowane dostarczenie wartości z AI w warunkach zgodności regulacyjnej, audytowalności i pomiaru jakości rozumowania systemu

WDROŻENIE I ZARZĄDZANIE

CogOps *Cognitive Operations — operacje kognitywne*

Operacyjna faza utrzymania systemów AI po wdrożeniu: monitoring, optymalizacja modeli, utrzymanie zgodności, dotrenowywanie, zarządzanie dryfem, świeżością wiedzy i degradacją jakości. Odpowiednik dojrzałego modelu operacyjnego (jak ITIL dla IT) dla systemów kognitywnych

PODSTAWY AI

Compliance Gap Analysis *Analiza luk zgodności*

Realizowany w Fazie 0 proces diagnozujący braki organizacji względem wymogów EU AI Act, norm ISO, regulacji sektorowych oraz wzorca umownego Ministerstwa Cyfryzacji. Stanowi podstawę planu naprawczego i dalszych decyzji architektonicznych.

METODYKA CDF

Compliance Readiness Report

Kluczowy artefakt F1.5 dokumentujący stan gotowości regulacyjnej organizacji.

METODYKA CDF

Compliance Timeline Tracker

Interaktywny widok deadline'ów regulacyjnych (EU AI Act, CRA, NIS2/uKSC, CEN/CENELEC) zintegrowany z CDF Platform. Umożliwia śledzenie terminów zgodności i automatyczne powiadomienia.

METODYKA CDF

Compliance-First *Compliance-first — zgodność jako warunek wstępny*

Podejście, w którym wymagania prawne, regulacyjne i audytowe traktuje się jako warunek wstępny wdrożenia AI, a nie etap końcowy. Zgodność projektuje się od początku i domyka przed odbiorem rozwiązania — co jest tańsze niż jej późniejsze „doklejanie”

WDROŻENIE I ZARZĄDZANIE

Confidence Calibration *Kalibracja pewności*

Miara tego, jak deklarowana przez model pewność odpowiedzi odpowiada jej faktycznej poprawności. Dobrze skalibrowany model, który mówi „jestem pewien w 90%”, ma rację w około 90% takich przypadków — dzięki temu człowiek wie, kiedy ufać, a kiedy sprawdzać

PODSTAWY AI

Conformity Assessment (Ocena zgodności) *Ocena zgodności systemu AI*

Formalna weryfikacja, czy system AI wysokiego ryzyka spełnia wymagania AI Act, przeprowadzana przed wprowadzeniem go do obrotu lub oddaniem do użytku (art. 43). Dla większości systemów z załącznika III jest to ocena wewnętrzna dostawcy; dla systemów biometrycznych i wyrobów objętych osobnymi przepisami — ocena z udziałem jednostki notyfikowanej. Obowiązki te stosuje się od 02.12.2027 (załącznik III) i 02.08.2028 (załącznik I)

REGULACJE I NORMY · W regulacji: AI Act

Coordination Effectiveness *Skuteczność koordynacji*

Metryka warstwy Agent SLA (Tabela 52 CDF) mierząca jakość współpracy między agentami w realizacji złożonych zadań end-to-end.

METODYKA CDF

CORE (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE zawsze aktywny – rdzeń CDF obowiązkowy dla każdego wdrożenia AI niezależnie od profilu organizacji. Obejmuje podstawowe fazy, kontrole i artefakty.

METODYKA CDF

CRA (Cyber Resilience Act) *Cyber Resilience Act — rozporządzenie (UE) 2024/2847*

Rozporządzenie (UE) 2024/2847 ustanawiające obowiązkowe wymagania cyberbezpieczeństwa dla produktów z elementami cyfrowymi — od sprzętu po oprogramowanie. Obejmuje bezpieczeństwo już na etapie projektowania (wymagania zasadnicze z załącznika I), obowiązki producentów (art. 13), zgłaszanie aktywnie wykorzystywanych podatności w 24 h / 72 h / 14 dni (art. 14), spis komponentów oprogramowania (SBOM), klasyfikację produktów ważnych i krytycznych oraz oznakowanie CE. Zgłaszanie podatności obowiązuje od 11.09.2026, pełne stosowanie od 11.12.2027. Kary: do 15 mln EUR lub 2,5% rocznego obrotu

REGULACJE I NORMY

Critical ICT Provider (Krytyczny dostawca ICT) *Krytyczny zewnętrzny dostawca usług ICT (DORA)*

Zewnętrzny dostawca usług ICT wyznaczony przez Europejskie Urzędy Nadzoru (EBA, ESMA, EIOPA) jako krytyczny dla sektora finansowego UE (DORA, art. 31–44). Podlega bezpośredniemu nadzorowi wiodącego organu, który może żądać informacji, prowadzić kontrole i wydawać zalecenia — także wobec dostawców chmury i modeli AI, z których korzystają banki i ubezpieczyciele

REGULACJE I NORMY · W regulacji: DORA

Cross-Validation Guardrail *Reguła walidacji krzyżowej ACE*

Automatyczna reguła sprawdzająca spójność konfiguracji ACE. 7 guardrails (CV-01 do CV-07) weryfikują np. wymuszenie REGULATED przy FINANCE, AGENT przy multi-agent, HCG-Full przy REGULATED, Tier 3 przy SCALE. Stosują trzy typy reguł: AUTO-CORRECT (automatyczna aktywacja brakującego tagu), BLOCKING (konfiguracja niedozwolona) oraz ADVISORY (ostrzeżenie wymagające świadomej decyzji). Reguły CV-01 do CV-04: AUTO-CORRECT. CV-05 do CV-07: ADVISORY (sygnalizacja konsekwencji). MV-01 do MV-07: ADVISORY (ostrzeżenia merytoryczne).

METODYKA CDF

CV-12 *Cross-Validation Rule 12 – ograniczenie LOA w procesach rozstrzygających*

Reguła walidacji krzyżowej ACE typu BLOCKING, wprowadzona w 1.6.0. Deklaracja procesu rozstrzygającego w podpytaniu 3a Kwestionariusza Konfiguracyjnego wraz z tagiem GOV blokuje przypisanie poziomu autonomii powyżej L2. Warstwa 0 walidacji obejmuje od tego wydania dwie reguły: CV-00 i CV-12.

METODYKA CDF

D

Data Act (rozp. 2023/2854) *Akt w sprawie danych — rozporządzenie (UE) 2023/2854*

Rozporządzenie (UE) 2023/2854 o zharmonizowanych przepisach dotyczących sprawiedliwego dostępu do danych i ich wykorzystywania. Daje użytkownikom prawo dostępu do danych generowanych przez ich urządzenia i usługi, a klientom usług chmurowych — prawo do zmiany dostawcy bez barier technicznych i umownych (przeciwdziała uzależnieniu od dostawcy). Stosowane od 12 września 2025 r.

REGULACJE I NORMY

Data Gravity *Grawitacja danych (data gravity)*

Tendencja danych do „przyciągania” aplikacji i usług: im więcej danych zgromadzono w jednym miejscu, tym trudniej i drożej je przenieść i tym bardziej determinują one architekturę docelową. Kluczowe przy decyzji, gdzie ma działać model AI — bliżej danych, a nie odwrotnie

SUWERENNOŚĆ I ARCHITEKTURA

Data Gravity Assessment *Analiza grawitacji danych*

Proces Fazy 0 badający cechy i położenie danych pod kątem ich wpływu na architekturę wdrożenia AI. Ocenia wolumen, tempo zmian, wymagania integracyjne, opóźnienia, regulacje, koszty zmiany i cykl życia danych.

METODYKA CDF

Data Gravity Index (DGI) *Indeks grawitacji danych*

Siedmiowymiarowa miara grawitacji danych obliczana w Fazie 0 na podstawie wymiarów: Volume, Velocity, Latency, Integration, Regulation, Switching Cost i Lifecycle. Wynik pomaga określić, czy najlepsza będzie architektura sovereign on-premise, hybryd sovereign-cloud czy air-gapped defense.

METODYKA CDF

Data Provenance Tracking *Śledzenie pochodzenia danych (data provenance)*

Mechanizm rejestrowania pełnej drogi danych — od źródła, przez przekształcenia, po wykorzystanie w odpowiedzi lub decyzji agenta AI. Zapewnia audytowalność, wspiera obowiązki z AI Act i pozwala odtworzyć, na jakich danych oparto każdą decyzję systemu

SUWERENNOŚĆ I ARCHITEKTURA · W regulacji: AI Act

DEFENSE (tag ACE) *DEFENSE*

Tag ACE aktywowany dla sektora obronnego (MON, Siły Zbrojne, infrastruktura krytyczna państwa). Włącza sekcje NATO PRU, K12 CBP, F6-Full, retraining pipeline. Zawsze aktywowany razem z tagami SOVEREIGN i GOV.

METODYKA CDF

Degradacja (modelu) *Model Degradation*

Spadek skuteczności i poprawności działania modelu AI w środowisku produkcyjnym — niezależnie od przyczyny (dryf danych, zmiana procesu, błędy integracji). Wymaga stałego pomiaru i traktowania na równi z ryzykiem operacyjnym

PODSTAWY AI

Delegation Chain Protocol *Delegation Chain Protocol (CDF, na bazie NIST NCCoE)*

Protokół łańcucha delegacji między agentami AI. Konstrukcja własna CDF rozwijająca filar autoryzacji z concept paper NIST NCCoE (5.02.2026). Nie jest cytatem z publikacji NIST. Kontrola T-012 ext., faza F2, artefakt Delegation Chain Log.

METODYKA CDF

Deployer (Podmiot wdrażający) *Podmiot stosujący system AI (deployer)*

Osoba fizyczna lub prawna, organ publiczny albo inny podmiot, który korzysta z systemu AI pod własną kontrolą w ramach działalności zawodowej (art. 3 pkt 4 AI Act; w polskim tekście „podmiot stosujący”). Wobec systemów wysokiego ryzyka ma obowiązki z art. 26: stosowanie zgodnie z instrukcją, nadzór człowieka, monitorowanie, przechowywanie logów i informowanie osób, których dotyczą decyzje, a w sektorze publicznym — także ocenę skutków dla praw podstawowych (art. 27)

REGULACJE I NORMY · W regulacji: AI Act

DevSecOps *DevSecOps — Development, Security and Operations*

Podejście integrujące praktyki bezpieczeństwa z procesami wytwarzania i utrzymania oprogramowania: testy bezpieczeństwa, skanowanie zależności i kontrola konfiguracji są częścią potoku CI/CD, a nie osobnym etapem. W projektach AI obejmuje też modele, prompty i komponenty agentowe

SUWERENNOŚĆ I ARCHITEKTURA

DGA *Akt w sprawie zarządzania danymi — rozporządzenie (UE) 2022/868*

Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2022/868 w sprawie europejskiego zarządzania danymi. Ustanawia ramy ponownego wykorzystywania chronionych danych sektora publicznego, działalności usług pośrednictwa danych oraz altruizmu danych; ma znaczenie przy projektach AI w administracji korzystających z danych publicznych

REGULACJE I NORMY

DGI *Data Gravity Index*

Skrót od Data Gravity Index. Zob. Data Gravity Index (DGI).

METODYKA CDF

Digital Omnibus (EU) *Digital Omnibus on AI — rozporządzenie (UE) 2026/1744*

Rozporządzenie (UE) 2026/1744 zmieniające AI Act oraz rozporządzenia 2018/1139 i 2023/1230 w celu uproszczenia wdrażania przepisów o sztucznej inteligencji. Najważniejsze zmiany: obowiązki dla samodzielnych systemów wysokiego ryzyka (załącznik III) stosowane od 02.12.2027 zamiast 02.08.2026, dla systemów wbudowanych w produkty (załącznik I) od 02.08.2028; okres przejściowy na maszynowo odczytywalne znakowanie treści syntetycznych do 02.12.2026; termin na krajowe piaskownice regulacyjne do 02.08.2027; nowe zakazy w art. 5 (intymne wizerunki bez zgody, materiały CSAM) stosowane od 02.12.2026; złagodzony obowiązek kompetencji w zakresie AI. Opublikowane w Dz.Urz. UE 24.07.2026, w mocy od 27.07.2026

REGULACJE I NORMY · W regulacji: AI Act · rozp. 2026/1744

Digital Resilience Testing (TLPT) *Testowanie cyfrowej odporności operacyjnej (DORA, TLPT)*

Program testowania odporności operacyjnej wymagany przez DORA od podmiotów finansowych (art. 24–27): coroczne testy systemów ICT oraz — dla wyznaczonych podmiotów — zaawansowane testy penetracyjne oparte na scenariuszach realnych zagrożeń (TLPT), co najmniej raz na trzy lata, z udziałem dostawców ICT wspierających krytyczne funkcje

REGULACJE I NORMY · W regulacji: DORA

DoD (Definition of Done) *Definicja ukończenia (Definition of Done)*

Zbiór formalnych kryteriów, które muszą być spełnione, aby etap pracy uznać za zakończony. W projektach AI obejmuje nie tylko działające rozwiązanie, lecz także artefakty zgodności, testy jakości odpowiedzi i zatwierdzenie przez osobę odpowiedzialną; dopiero wtedy można przejść do kolejnej fazy

WDROŻENIE I ZARZĄDZANIE

DORA *Digital Operational Resilience Act — rozporządzenie (UE) 2022/2554*

Rozporządzenie (UE) 2022/2554 w sprawie operacyjnej odporności cyfrowej sektora finansowego, stosowane od 17 stycznia 2025 r. Obejmuje zarządzanie ryzykiem ICT, zgłaszanie poważnych incydentów, testowanie odporności, zarządzanie ryzykiem ze strony zewnętrznych dostawców ICT (w tym dostawców chmury i AI) oraz wymianę informacji o zagrożeniach. Wymaga m.in. klauzul umownych, strategii wyjścia i rejestru umów z dostawcami ICT

REGULACJE I NORMY

DPIA (Data Protection Impact Assessment) *Ocena skutków dla ochrony danych (DPIA)*

Ocena skutków planowanego przetwarzania danych osobowych dla praw i wolności osób fizycznych, wymagana przez art. 35 RODO, gdy przetwarzanie może powodować wysokie ryzyko — w szczególności przy profilowaniu, zautomatyzowanym podejmowaniu decyzji i przetwarzaniu na dużą skalę, co obejmuje większość wdrożeń AI na danych pracowników lub klientów

REGULACJE I NORMY · W regulacji: RODO

DPO / IOD (Inspektor Ochrony Danych) *Inspektor ochrony danych (IOD, ang. DPO)*

Osoba wyznaczona na podstawie art. 37 RODO, odpowiedzialna za monitorowanie zgodności przetwarzania danych osobowych, doradztwo (w tym przy ocenie skutków dla ochrony danych) i współpracę z organem nadzorczym (art. 39). Przy wdrożeniach AI powinna być włączana na etapie projektowania, a nie po uruchomieniu

REGULACJE I NORMY · W regulacji: RODO

Drift (Model Drift) *Dryf modelu*

Stopniowe pogarszanie się jakości odpowiedzi modelu w czasie, gdy dane, na których pracuje, zaczynają odbiegać od tych, na których był trenowany lub kalibrowany — nowe przepisy, nowe produkty, zmiana języka klientów. Wykrywa się go monitoringiem po wdrożeniu, a leczy dotrenowaniem lub odświeżeniem bazy wiedzy

PODSTAWY AI

Dual-Orbit Sovereignty Matrix

Narzędzie oceny ryzyka geopolitycznego stacku AI. Identyfikuje trzy orbity dostawców: US-orbit, CN-orbit, EU-sovereign. Ocenia: dostępność modeli, ryzyko sankcji, data residency, supply chain resilience, IP protection. Patrz: Załącznik F, sekcja 2.9.

METODYKA CDF

E

ECS *Enterprise Cognitive Systems — systemy kognitywne klasy korporacyjnej*

Systemy kognitywne klasy korporacyjnej — platformy AI zintegrowane z procesami organizacji, które wspierają decyzje, automatyzację, analizę wiedzy i pracę agentów AI. W odróżnieniu od pojedynczego narzędzia AI obejmują dane, wiedzę, nadzór i zgodność jako jedną całość

PODSTAWY AI

EHDS (rozp. 2025/327) *Europejska przestrzeń danych dotyczących zdrowia — rozporządzenie (UE) 2025/327*

Rozporządzenie (UE) 2025/327 ustanawiające europejską przestrzeń danych dotyczących zdrowia. Reguluje pierwotne wykorzystanie danych (dostęp pacjentów, interoperacyjność elektronicznej dokumentacji) i wtórne — do badań, innowacji i trenowania AI — na podstawie zezwoleń wydawanych przez organy ds. dostępu do danych. Stosowane etapowo od 26 marca 2027 r., pełna funkcjonalność w latach 2029–2031

REGULACJE I NORMY

EN 18286:2026 *EN 18286:2026 — system zarządzania jakością AI na potrzeby AI Act*

Pierwsza europejska norma zharmonizowana opracowana na potrzeby AI Act, opublikowana przez CEN/CENELEC w 2026 r. Opisuje system zarządzania jakością, którego art. 17 AI Act wymaga od dostawców systemów wysokiego ryzyka. Do czasu przywołania w Dzienniku Urzędowym UE stosowanie normy nie daje domniemania zgodności, ale porządkuje wymagania i ułatwia audyt

Enterprise AI Catalog *Katalog rozwiązań AI w organizacji*

Wewnętrzna biblioteka gotowych i zweryfikowanych przypadków użycia AI, dostępnych do ponownego wdrożenia w różnych częściach organizacji — z opisem zastosowania, wymaganych danych, ryzyk i statusu zgodności

WDROŻENIE I ZARZĄDZANIE

EU Database (Unijna baza danych systemów AI) *Unijna baza danych systemów AI wysokiego ryzyka*

Baza prowadzona przez Komisję Europejską (art. 71 AI Act), w której dostawcy rejestrują systemy AI wysokiego ryzyka przed wprowadzeniem do obrotu, a podmioty publiczne — przed rozpoczęciem ich używania (art. 49). Większość wpisów jest publicznie dostępna; dane systemów stosowanych przez organy ścigania i w obszarze migracji widzą tylko organy nadzoru

REGULACJE I NORMY · W regulacji: AI Act

Evidence Book *Evidence Book — warstwa dowodowa metodyki*

Zbiór materiałów potwierdzających skuteczność i zgodność metodyki: mapowanie na regulacje, studia przypadków, przykładowe produkty wdrożenia, warunki brzegowe składanych deklaracji oraz benchmarki i metryki. Podstawa, na którą organizacja może powołać się wobec audytora lub regulatora

WDROŻENIE I ZARZĄDZANIE

Executive Alignment Workshop *Warsztat uzgodnienia kierownictwa*

Warsztat dla najwyższego kierownictwa na starcie programu AI, którego celem jest uzgodnienie priorytetów biznesowych, oczekiwanych rezultatów, akceptowalnego ryzyka i horyzontu czasowego. Bez niego inicjatywy AI rozpraszają się między działami i nie mają sponsora

WDROŻENIE I ZARZĄDZANIE

Exit Strategy *Strategia wyjścia (exit strategy)*

Plan planowego lub awaryjnego zakończenia współpracy z dostawcą AI, obejmujący przeniesienie danych, kodu, konfiguracji, dokumentacji i wiedzy operacyjnej do innego dostawcy lub do własnej infrastruktury. Wymagana przez DORA dla usług wspierających funkcje krytyczne i zalecana przy każdym wdrożeniu opartym na zewnętrznym modelu

WDROŻENIE I ZARZĄDZANIE · W regulacji: DORA

F

F6-AFC-01 *Specyfikacja kanału zgłaszania obaw AI*

Artefakt CDF definiowany w Fazie F6 (Operate): AI Concerns Feedback Channel Spec. Dokument określający: kanały przyjęcia zgłoszeń (formularz, e-mail, dedykowany portal), procedurę obsługi, SLA odpowiedzi, ścieżkę eskalacji oraz reguły archiwizacji. Spełnia wymogi art. 9 ust. 9 AI Act dotyczące zapewnienia osobom, których dotyczy system AI, możliwości składania skarg i wniosków.

METODYKA CDF

F6-Full *Pełny tryb Fazy 6*

Rozszerzony tryb Fazy 6 (CogOps) aktywowany tagiem SCALE: pełne dashboardy CogOps, Knowledge Freshness Index, automatyczne retraining triggers, Cognitive Degradation Monitoring, Capacity Planning.

METODYKA CDF

F6-Lite *Uproszczony tryb Fazy 6*

Podstawowy tryb Fazy 6 (CogOps) aktywowany tagiem PRODUCTION: monitoring produkcyjny, Cognitive SLA Tier 1, alerty degradacji, bazowy retraining schedule. Wystarczający dla wdrożeń produkcyjnych bez pełnej skali.

METODYKA CDF

FINANCE (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla sektora finansowego. Włącza sekcje DORA, ICT third-party risk, specyfikę KNF. Automatycznie wymusza REGULATED (guardrail CV-03).

METODYKA CDF

Fine-tuning *Dostrajanie modelu*

Dalsze trenowanie gotowego modelu na wybranym zbiorze danych, żeby dostosować go do domeny, stylu lub typu zadań. W odróżnieniu od RAG zmienia sam model, nie tylko źródła, z których korzysta — dlatego jest droższe i trudniejsze do audytu

PODSTAWY AI

FRIA *Ocena skutków dla praw podstawowych (FRIA)*

Ocena skutków dla praw podstawowych wymagana przez art. 27 AI Act przed pierwszym użyciem systemu AI wysokiego ryzyka przez podmioty prawa publicznego, podmioty świadczące usługi publiczne oraz instytucje oceniające zdolność kredytową lub ryzyko ubezpieczeniowe osób fizycznych. Opisuje procesy, osoby dotknięte, ryzyka i środki nadzoru; wynik zgłasza się organowi nadzoru rynku

REGULACJE I NORMY · W regulacji: AI Act

Fully On-Premise *Pełne wdrożenie lokalne (fully on-premise)*

Wzorzec wdrożenia, w którym cała warstwa AI — modele, dane, logi, orkiestracja agentów i integracje — pozostaje w infrastrukturze organizacji. Stosowany tam, gdzie decydujące są suwerenność danych, kontrola operacyjna i ograniczenie ryzyka związanego z dostawcą

SUWERENNOŚĆ I ARCHITEKTURA

G

GDPR *General Data Protection Regulation — rozporządzenie (UE) 2016/679*

Ogólne rozporządzenie o ochronie danych (UE) 2016/679, stosowane bezpośrednio we wszystkich państwach UE od 25 maja 2018 r. W Polsce funkcjonuje pod nazwą RODO; krajowa ustawa o ochronie danych osobowych jedynie je uzupełnia. Zob. RODO

REGULACJE I NORMY · W regulacji: RODO

GEO (tag ACE) *Tag lokalizacyjny ACE*

Warstwa jurysdykcyjna ACE działająca niezależnie od tagów sektorowych. Warianty: GEO:PL (Polska), GEO:EU (inny kraj UE/EOG), GEO:MULTI (wielojurysdykcyjne), GEO:NON-EU (poza UE).

METODYKA CDF

Global Rollout *Globalne wdrożenie / skalowanie*

Proces realizowany w Fazie 5, polegający na replikacji ustandaryzowanych agentów, wzorców i komponentów AI na kolejne działy, jednostki lub obszary działalności.

METODYKA CDF

Gotowość dowodowa *Gotowość dowodowa (evidence readiness)*

Zdolność organizacji do przedstawienia — na wypadek roszczenia z tytułu odpowiedzialności za produkt (dyrektywa 2024/2853) — kompletnego i wiarygodnego materiału dowodowego: rejestru agentów, logów decyzji, dokumentacji modelu i testów. Taki materiał pozwala wykazać należytą staranność i obalić domniemania wadliwości oraz związku przyczynowego z art. 10 dyrektywy

WDROŻENIE I ZARZĄDZANIE · W regulacji: PLD 2024/2853

GOV (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla sektora publicznego lub obronnego. Włącza sekcje specyficzne dla administracji, NATO PRU, wzorzec MC, informacje niejawne. Automatycznie wymusza REGULATED.

METODYKA CDF

Governance Divergence Matrix *Matryca rozbieżności governance US-EU*

Narzędzie analityczne CDF porównujące podejścia regulacyjne USA i UE w obszarze AI governance. Obejmuje 8 wymiarów: Risk Classification, Technical Documentation, Human Oversight, Transparency, Post-Market Monitoring, Liability, Conformity Assessment, Cross-Border Transfer. Dla każdego wymiaru wskazuje podejście EU (AI Act) vs US (NIST AI RMF + EO 14110), poziom rozbieżności i rekomendacje CDF dla organizacji działających w obu jurysdykcjach. Patrz: sekcja K5/K11 w Metodyce.

METODYKA CDF

Ground Truth *Prawda referencyjna*

Zweryfikowany, poprawny stan wiedzy lub wzorcowa odpowiedź, które służą jako punkt odniesienia przy pomiarze trafności AI. Bez zestawu takich odpowiedzi nie da się rzetelnie zmierzyć poziomu halucynacji

H

Hallucination *Halucynacja*

Odpowiedź modelu, która brzmi wiarygodnie, ale jest nieprawdziwa, niepotwierdzona albo nie ma oparcia w źródłach — zmyślony przepis, nieistniejący wyrok, błędna liczba. Ryzyko ogranicza się przez oparcie odpowiedzi na kontrolowanej bazie wiedzy (RAG), wymaganie cytowania źródeł i ciągły pomiar

PODSTAWY AI

Hallucination Rate *Wskaźnik halucynacji*

Odsetek odpowiedzi systemu AI, które zawierają informacje zmyślone lub sprzeczne z bazą wiedzy — mierzony na zestawie testowym przed wdrożeniem i cyklicznie po nim. Dla procesów krytycznych przyjmuje się progi na poziomie pojedynczych procent

PODSTAWY AI

HCG *HCG — Human Competence Gate*

Skrót od Human Competence Gate — mechanizmu sprawdzającego, czy użytkownik rozumie kluczowe fakty i konsekwencje decyzji AI, zanim ją zatwierdzi. Zob. Human Competence Gate

WDROŻENIE I ZARZĄDZANIE

HCG Audit Log Specification

Artefakt F2 definiujący format logowania interakcji człowiek-AI w ramach Human Competence Gate.

METODYKA CDF

HCG Control Policy

Artefakt F2 definiujący zasady nadzoru ludzkiego nad systemem AI: kto zatwierdza jakie decyzje, przy jakim progu pewności.

METODYKA CDF

HCG-Full *Pełny tryb Human Competence Gate*

Rozszerzony tryb HCG aktywowany tagami AGENT + REGULATED: pełna matryca HCG, checkpoint per use case, audit log decyzji, rejestr pytań kontrolnych, eskalacja, periodic review.

METODYKA CDF

HCG-Lite *Uproszczony tryb Human Competence Gate*

Podstawowy tryb HCG aktywowany tagiem CORE: minimum jeden checkpoint ludzki per wdrożenie AI, decyzja go/no-go przed uruchomieniem produkcyjnym, jasna definicja kto odpowiada za nadzór.

METODYKA CDF

HDC *Health-AI Compliance*

Mechanizm zgodności AI w sektorze zdrowia, osadzający w metodyce reżim aktywowany wejściem na rynek medyczny: EHDS (dostęp i wtórne wykorzystanie danych zdrowotnych), MDR/IVDR (kwalifikacja systemu AI jako oprogramowania wyrobu medycznego, sprzężona ze ścieżką wysokiego ryzyka zał. I AI Act) oraz RODO art. 9 (dane szczególnej kategorii). Bramkowany tagiem HEALTH. Obejmuje kontrole HDC-1 i HDC-2. Fazy F1.5/F2. Artefakty: A-HDP-01, A-MDR-01. Wprowadzony w CDF 1.5.0.

METODYKA CDF

HDC-1 *Health Data Governance*

Kontrola CDF mechanizmu HDC: przetwarzanie danych szczególnej kategorii (RODO art. 9), wtórne wykorzystanie i interoperacyjność zgodnie z EHDS. Rozszerzenie mechanizmu K3 Data Governance (cegła istnieje – nie nowy mechanizm). Obszar wiedzy KD3. Fazy F1.5/F2. Artefakt A-HDP-01. Wprowadzona w CDF 1.5.0.

METODYKA CDF

HDC-2 *Medical Device AI Conformity*

Kontrola CDF mechanizmu HDC: ocena kwalifikacji systemu AI jako oprogramowania wyrobu medycznego (MDR 2017/745 / IVDR 2017/746) i sprzężenie ze ścieżką wysokiego ryzyka zał. I AI Act. Nowa kontrola – CDF nie posiadał dotąd zgodności wyrobu medycznego. Obszar wiedzy KD8. Faza F1.5. Artefakt A-MDR-01. Wprowadzona w CDF 1.5.0.

METODYKA CDF

HEALTH (tag ACE) *HEALTH - tag konfiguracyjny ACE*

Tag konfiguracyjny ACE aktywowany, gdy wdrożenie dotyczy sektora zdrowia (dane zdrowotne, wyrób medyczny AI, EHDS). Włącza mechanizm HDC (kontrolę HDC-1/HDC-2) oraz reguły walidacji krzyżowej CV-10 (auto-aktywacja tagu REGULATED) i CV-11. Powiązany z profilem referencyjnym HEALTH. Wprowadzony w CDF 1.5.0.

METODYKA CDF

HITL *Human-in-the-Loop — człowiek w pętli decyzyjnej*

Model nadzoru, w którym człowiek jest częścią procesu decyzyjnego AI: zatwierdza, koryguje lub blokuje decyzję, zanim wywoła ona skutek. Jeden z podstawowych sposobów realizacji wymogu nadzoru człowieka z art. 14 AI Act

PODSTAWY AI · W regulacji: AI Act

Human Competence Gate (HCG) *Human Competence Gate — kontrola kompetencji przed zatwierdzeniem decyzji AI*

Mechanizm, w którym system AI — zanim dopuści użytkownika do zatwierdzenia decyzji — generuje zależny od kontekstu zestaw 5-15 pytań tak/nie sprawdzających zrozumienie istotnych faktów, konsekwencji, ograniczeń i wyjątków w danej sprawie. Nieosiągnięcie progu poprawności blokuje zatwierdzenie i może skutkować eskalacją lub skierowaniem do materiału szkoleniowego. Praktyczne wzmocnienie nadzoru człowieka w procesach podwyższonego ryzyka

WDROŻENIE I ZARZĄDZANIE

Human Factors Monitoring *Monitorowanie czynnika ludzkiego (human factors monitoring)*

Kategoria monitorowania systemu AI po wdrożeniu (wg NIST AI 800-4) obejmująca wpływ systemu na ludzi, którzy z nim pracują: rozbieżność między intencją a tym, co użytkownik widzi, „ciemne wzorce” interakcji (schlebianie, antropomorfizacja), pętla sprzężenia zwrotnego. W praktyce wdrożeniowej rozszerza się ją o bezrefleksyjne akceptowanie podpowiedzi (automation bias), zanik umiejętności i nadmierne zaufanie. Wymaga obserwacji zachowań użytkowników, nie tylko metryk modelu

WDROŻENIE I ZARZĄDZANIE

Human Override *Interwencja człowieka*

Możliwość zatrzymania, zablokowania, cofnięcia albo zmiany decyzji podjętej przez system AI przez upoważnioną osobę — w każdym momencie, bez zgody systemu. W sektorach regulowanych to podstawowy element realnego, a nie tylko formalnego nadzoru

PODSTAWY AI

Hybrid Sovereign-Cloud *Hybrydowa architektura suwerenna*

Wzorzec wdrożenia, w którym dane krytyczne, pamięć organizacyjna, kontrola dostępu i kluczowe mechanizmy nadzoru pozostają pod kontrolą organizacji, a część mocy obliczeniowej lub modeli działa w zaufanym środowisku chmurowym — najczęściej w jurysdykcji UE, z gwarancjami umownymi i technicznymi

SUWERENNOŚĆ I ARCHITEKTURA

I

ICT Security Addendum *Aneks bezpieczeństwa ICT*

Uzupełnienie kontraktowe stosowane zwłaszcza w kontekście NIS2/uKSC, obejmujące klauzule bezpieczeństwa, prawo audytu, zarządzanie incydentami, backup, strategię wyjścia i ocenę dostawców ICT oraz AI. Wdrażany w Fazie 1.5 CDF.

METODYKA CDF

Immutable Audit Trail *Niemodyfikowalny dziennik zdarzeń (immutable audit trail)*

Dziennik zdarzeń, którego wpisów nie można usunąć ani zmienić bez pozostawienia śladu — dzięki podpisom kryptograficznym, łańcuchowaniu wpisów lub nośnikom tylko do zapisu.

Podstawa audytowalności i dowodowości działań systemów AI i agentów, wymagana m.in. przez AI Act (rejestrowanie zdarzeń)

SUWERENNOŚĆ I ARCHITEKTURA · W regulacji: AI Act

Inference / Inferencja *Wnioskowanie modelu*

Faza działania modelu AI, w której na podstawie nowych danych wejściowych generuje on odpowiedzi, przewidywania lub rekomendacje — w odróżnieniu od trenowania. To inferencja, nie trening, generuje koszty w codziennej eksploatacji

PODSTAWY AI

Innovation Plateau Diagnostic *Diagnostyka plateau innowacji*

Narzędzie diagnozujące stan, w którym inwestycje w AI przestają przekładać się na wartość biznesową. Rozróżnia pięć typów blokad: czyścić pilotaży (pilot purgatory), paraliż skalowania, zator decyzyjny w nadzorze (governance gridlock), opór kulturowy i niegotowość danych. Każdy typ ma odrębny plan naprawczy

WDROŻENIE I ZARZĄDZANIE

ISO/IEC 23894 *ISO/IEC 23894:2023 — zarządzanie ryzykiem w AI*

Norma zawierająca wytyczne zarządzania ryzykiem związanym ze sztuczną inteligencją — identyfikacji, analizy, oceny i ograniczania ryzyk na etapie projektowania, wdrażania i użytkowania systemów AI. Rozwija ogólne zasady ISO 31000 o specyfikę AI i jest naturalnym uzupełnieniem ISO/IEC 42001

REGULACJE I NORMY · W regulacji: ISO/IEC 42001

ISO/IEC 42001 *ISO/IEC 42001:2023 — system zarządzania AI*

Pierwsza międzynarodowa norma certyfikowalna dla systemów zarządzania sztuczną inteligencją (AIMS). Określa wymagania dotyczące polityki AI, ról, oceny ryzyka i wpływu, cyklu życia systemów i ciągłego doskonalenia; załącznik A zawiera katalog kontroli. Adresowana do organizacji, które rozwijają lub wdrażają AI i chcą to wykazać wobec klientów, audytorów i regulatora

REGULACJE I NORMY

IVDR (rozp. 2017/746) *Rozporządzenie w sprawie wyrobów medycznych do diagnostyki in vitro — (UE) 2017/746*

Rozporządzenie (UE) 2017/746 w sprawie wyrobów medycznych do diagnostyki in vitro. Obejmuje także oprogramowanie diagnostyczne, w tym algorytmy AI analizujące wyniki badań laboratoryjnych; takie systemy podlegają jednocześnie IVDR i przepisom AI Act o systemach wysokiego ryzyka

REGULACJE I NORMY · W regulacji: AI Act · MDR/IVDR

J

J-curve of Productivity *Krzywa J produktywności*

Zjawisko przejściowego spadku produktywności na początku wdrożenia AI — gdy pracownicy uczą się nowych narzędzi, procesy są przebudowywane, a baza wiedzy dopiero powstaje — po którym następuje przyspieszenie przewyższające stan wyjściowy. Wdrożenie powinno ten spadek przewidywać i planować, zamiast traktować go jako porażkę

WDROŻENIE I ZARZĄDZANIE

Jednostka notyfikowana (SSI) *Jednostka notyfikowana (AI Act, ustawa o systemach SI)*

Jednostka oceniająca zgodność, notyfikowana zgodnie z rozdziałem III sekcją 4 AI Act, uprawniona do oceny systemów AI wysokiego ryzyka z udziałem strony trzeciej. W Polsce organem notyfikującym jest minister właściwy do spraw informatyzacji, a akredytację jednostek prowadzi Polskie Centrum Akredytacji (ustawa o systemach sztucznej inteligencji)

REGULACJE I NORMY · W regulacji: AI Act · ustawa o systemach SI

Jurisdiction Compliance Card *Karta jurysdykcyjna*

Artefakt ACE generowany w F0: dokument opisujący prawo bazowe, prawo krajowe, instytucje nadzorcze, timeline wymogów i artefakty specyficzne dla lokalizacji wdrożenia.

METODYKA CDF

K

Katalog artefaktów CDF *CDF Artifact Catalogue*

Jedyne źródło wykazu artefaktów CDF, wprowadzone w 1.6.0. Scala wykazy dotąd rozproszone w listach fazowych, artefaktach ACE, artefaktach CBP, rejestrze artefaktów SSI i opisach mechanizmów dziedzicznych. 94 pozycje z przypisaniem fazy, tagu aktywacji, właściciela z macierzy RACI i typu nośnika. Liczba artefaktów aktywnych dla wdrożenia wynika z profilu ACE, nie z liczby pozycji katalogu.

METODYKA CDF

Katalog kontroli CDF *CDF Control Catalogue*

Jedyne źródło wykazu kontroli CDF, wprowadzone w 1.6.0. 34 kontrole w ośmiu rodzinach: T (kontrole metodyczne), MCP, VDP, SAG, LER, CBC, HDC, PAG. Każda kontrola ma przypisany obszar wiedzy, fazy, tag aktywacji i powiązany artefakt.

KG-RAG Integration Governance *Zarządzanie integracją KG-RAG*

Zestaw zasad (Faza 6 CDF) zapewniających, że agenty generatywne korzystają wyłącznie z aktualnych, zwalidowanych źródeł w grafie wiedzy. Obejmuje mechanizmy Knowledge Freshness Monitoring i Knowledge Audit Trail.

METODYKA CDF

Kill/Pivot Criteria *Kryteria przerwania lub zmiany kierunku pilotażu*

Z góry ustalone warunki, przy których pilotaż AI należy zakończyć lub przekierować — np. brak oczekiwanej jakości odpowiedzi, koszt powyżej progu, brak adopcji. Odróżnić od bramy Scale-or-Kill: tam decyzja jest binarna (skaluj albo zakończ), a opcja „przedłużmy pilotaż” jest celowo wykluczona

WDROŻENIE I ZARZĄDZANIE

Knowledge Audit Trail *Ścieżka audytowa bazy wiedzy*

Pełna, niezmienna (immutable) ścieżka audytowa rejestrująca wszystkie zmiany w bazie wiedzy organizacji. Spełnia wymogi regulatorów w zakresie transparentności i możliwości kontroli systemu AI (Faza 6 CDF).

METODYKA CDF

Knowledge Curator *Kurator wiedzy*

Dedykowana rola organizacyjna odpowiedzialna za jakość, aktualność, bezpieczeństwo i wersjonowanie grafu wiedzy. Zarządza Ontology Change Control i wskaźnikami świeżości wiedzy. Rola wdrażana w Fazie 6 CDF.

METODYKA CDF

Knowledge Debt Management *Zarządzanie długiem wiedzy*

Systematyczne wykrywanie i usuwanie długu informacyjnego organizacji: braków w dokumentacji, silosów danych, nieaktualnych procedur i rozproszonej wiedzy. Nieuporządkowana wiedza jest najczęstszą przyczyną słabych odpowiedzi systemu AI

PODSTAWY AI

Knowledge Freshness Index *Indeks świeżości wiedzy*

Miara wykorzystywana w Cognitive SLA, określająca jaki odsetek bazy wiedzy został zaktualizowany w zdefiniowanym oknie czasowym. Metryka Cognitive SLA Tier 3 (aktywna od profilu SOVEREIGN/SCALE). Target CDF: $\geq 95\%$ w 7-dniowym oknie czasowym.

METODYKA CDF

Knowledge Freshness Monitoring

Automatyczny system alertów F6 wyzwalany przy przekroczeniu progu przedawnienia wiedzy w Knowledge Graph lub RAG.

METODYKA CDF

Knowledge Graph *Graf wiedzy*

Uporządkowana reprezentacja wiedzy organizacji w postaci obiektów (osób, dokumentów, procesów, pojęć) i relacji między nimi. Fundament systemów RAG: pozwala modelowi sięgać po fakty i powiązania zamiast zgadywać. Wymaga właściciela i kontroli zmian tak jak każdy rejestr

PODSTAWY AI

Komisja Rozwoju i Bezpieczeństwa Sztucznej Inteligencji *Komisja Rozwoju i Bezpieczeństwa Sztucznej Inteligencji (KRiBSI)*

Centralny organ nadzoru rynku sztucznej inteligencji w Polsce i pojedynczy punkt kontaktowy w rozumieniu AI Act, powołany ustawą z dnia 3 lipca 2026 r. o systemach sztucznej inteligencji. W jej skład wchodzi Przewodniczący, dwóch Zastępców oraz członkowie wskazani przez UOKiK, KNF, KRRiT i UKE. Komisja wydaje opinie indywidualne, prowadzi kontrole i postępowania, nakłada kary oraz prowadzi piaskownicę regulacyjną

REGULACJE I NORMY · W regulacji: AI Act

KPI *KPI — kluczowy wskaźnik efektywności*

Mierzalna miara stopnia realizacji celu biznesowego, operacyjnego lub wdrożeniowego. W projektach AI obok wskaźników biznesowych (czas obsługi, koszt sprawy) stosuje się wskaźniki jakości systemu, np. odsetek halucynacji lub trafność odpowiedzi

WDROŻENIE I ZARZĄDZANIE

Kryptografia postkwantowa (Post-Quantum Cryptography) *Kryptografia postkwantowa (PQC)*

Algorytmy kryptograficzne odporne na ataki z użyciem komputerów kwantowych, standaryzowane przez NIST od 2024 r. (FIPS 203–205). Organizacje powinny mieć plan migracji kryptograficznej — inwentaryzację użycia kryptografii i harmonogram wymiany — co przewiduje także Strategia Cyberbezpieczeństwa RP na lata 2025–2029

SUWERENNOŚĆ I ARCHITEKTURA

Kwestionariusz Konfiguracyjny *ACE Configuration Questionnaire*

Interaktywna ankieta ACE (7 pytań) służąca do ustalenia profilu konfiguracyjnego organizacji w trybie Self-Service. Pytania dotyczą: celu wdrożenia, typów AI, sektora, danych, skali, lokalizacji i compliance.

METODYKA CDF

L

Latency *Opóźnienie (czas odpowiedzi)*

Czas, jaki system potrzebuje na odpowiedź lub wykonanie operacji. W systemach AI ma znaczenie infrastrukturalne (wydajność modelu) i użytkowe (czy pracownik zaczeka na odpowiedź) — i bezpośrednio wpływa na koszt inferencji

PODSTAWY AI

LER *Liability Evidence Readiness*

Rodzina mechanizmów CDF zapewniająca gotowość dowodową pod odpowiedzialność za produkt wynikającą z dyrektywy 2024/2853 (PLD), która czyni oprogramowanie i systemy AI produktem objętym odpowiedzialnością niezależną od winy. LER spina istniejące nośniki dowodowe (Agent Registry, Immutable Audit Trail, Evidence Book) w reżim umożliwiający obalenie domniemań wadliwości i związku przyczynowego (art. 9–10 PLD) oraz realizację obowiązku ujawnienia dowodów. Obejmuje kontrole LER-1 i LER-2. Fazy F1.5/F2/F6. Artefakt: A-PLD-01. Wprowadzony w CDF 1.5.0.

METODYKA CDF

LER-1 *Product Liability Evidence Readiness*

Kontrola CDF mechanizmu LER: utrzymanie kompletu dowodowego (dokumentacja modelu i danych, logi decyzji, Agent Registry, wersjonowanie, testy bezpieczeństwa) w stanie umożliwiającym obalenie domniemań z art. 10 dyrektywy PLD 2024/2853 oraz realizację obowiązku ujawnienia z art. 9, z ochroną tajemnicy przedsiębiorstwa. Obszary wiedzy KD8/KD6/KD3. Fazy F1.5/F2/F6. Wprowadzona w CDF 1.5.0.

METODYKA CDF

LER-2 *Substantial Modification Register*

Kontrola CDF mechanizmu LER: rejestr znaczących modyfikacji (fine-tuning, ciągłe uczenie, aktualizacje) sygnalizujący moment, w którym podmiot wdrażający przechodzi w rolę producenta (art. 8 ust. 2 dyrektywy PLD 2024/2853) i przejmuje odpowiedzialność. Obszary wiedzy KD1/KD4. Fazy F2/F6. Wprowadzona w CDF 1.5.0.

METODYKA CDF

Liability Framework *Ramy odpowiedzialności za system AI*

Uporządkowany podział odpowiedzialności za decyzje i działania systemu AI w środowisku produkcyjnym: między dostawcą (provider), podmiotem stosującym (deployer) i wewnętrznym właścicielem systemu w organizacji. Punkt wyjścia stanowią obowiązki z art. 9 (zarządzanie ryzykiem) i art. 26 (obowiązki podmiotu stosującego) AI Act oraz przepisy o odpowiedzialności za produkt; ramy przekładają je na konkretne role, decyzje i zapisy umowne

REGULACJE I NORMY · W regulacji: AI Act

LLM *Large Language Model — duży model językowy*

Duży model językowy trenowany na bardzo dużych zbiorach tekstu, zdolny do generowania języka naturalnego, streszczania, klasyfikacji, ekstrakcji informacji i wykonywania poleceń agentów. Sam w sobie nie zna faktów organizacji — potrzebuje bazy wiedzy (RAG) lub dostrojenia

PODSTAWY AI

LOA *Level of Autonomy — poziom autonomii*

Poziom autonomii agenta AI w pięciostopniowej skali L0–L4: od agenta, który wyłącznie doradza, po agenta w pełni samodzielnego w zdefiniowanych ramach. Zobacz: Skala autonomii (L0–L4)

M

Make-vs-Buy Decision Framework *Decyzja „budować czy kupić” (make-vs-buy)*

Ustrukturyzowana ocena ekonomiczno-operacyjna uzasadniająca wybór między budową rozwiązania AI własnymi siłami a zakupem gotowych komponentów lub usług. Uwzględnia całkowity koszt posiadania, ryzyko, czas wdrożenia, kompetencje zespołu i uzależnienie od dostawcy

WDROŻENIE I ZARZĄDZANIE

Mapa Kognitywna *Cognitive Map*

Wizualizacja procesów, danych, kompetencji i obszarów decyzyjnych organizacji pod kątem potencjału wdrożenia AI — pokazuje, gdzie wiedza jest, gdzie jej brakuje i które decyzje da się wesprzeć systemem kognitywnym

PODSTAWY AI

MC *Ministerstwo Cyfryzacji*

Ministerstwo Cyfryzacji — urząd obsługujący ministra właściwego do spraw informatyzacji, odpowiedzialny m.in. za politykę AI i cyberbezpieczeństwa w Polsce, Krajowy System Cyberbezpieczeństwa, Przewodnik po AI dla administracji publicznej (marzec 2026) oraz projekt wzorcowych klauzul umownych dla zamówień publicznych na systemy AI. Minister cyfryzacji jest też organem notyfikującym w rozumieniu AI Act i zapewnia obsługę Komisji Rozwoju i Bezpieczeństwa SI

REGULACJE I NORMY · W regulacji: AI Act

MCP *Model Context Protocol*

Otwarty protokół (Anthropic, 2024) standaryzujący komunikację między agentami AI a zewnętrznymi narzędziami, źródłami danych i usługami. Ułatwia integrację, ale każde połączenie MCP jest nowym punktem wejścia do systemu — wymaga uwierzytelniania, limitów i audytu wywołań

PODSTAWY AI

MCP Security *Bezpieczeństwo protokołu MCP (Model Context Protocol)*

Zbiór wymagań bezpieczeństwa dla warstwy Model Context Protocol — otwartego standardu łączenia modeli AI z narzędziami i danymi. Obejmuje autoryzację opartą na OAuth 2.1 (specyfikacja MCP z 28 lipca 2026 r.), uwierzytelnianie serwerów i klientów, ograniczanie uprawnień narzędzi, ochronę przed wstrzykiwaniem poleceń przez treści z narzędzi oraz rejestrowanie wywołań. Źródła: specyfikacja MCP, dokument Security Best Practices oraz arkusz NSA „Model Context Protocol (MCP): Security Design Considerations for AI-Driven Automation” (maj 2026)

SUWERENNOŚĆ I ARCHITEKTURA

MDR (rozp. 2017/745) *Rozporządzenie w sprawie wyrobów medycznych — (UE) 2017/745*

Rozporządzenie (UE) 2017/745 w sprawie wyrobów medycznych. Obejmuje również oprogramowanie kwalifikowane jako wyrób medyczny, w tym systemy AI wspierające diagnostykę i decyzje terapeutyczne. Taki system podlega jednocześnie ocenie zgodności z MDR i — jako element bezpieczeństwa produktu — przepisom AI Act o systemach wysokiego ryzyka (załącznik I)

REGULACJE I NORMY · W regulacji: AI Act · MDR/IVDR

MFA *Uwierzytelnianie wieloskładnikowe (MFA)*

Mechanizm weryfikacji tożsamości wymagający co najmniej dwóch niezależnych czynników — np. hasła i klucza sprzętowego lub aplikacji uwierzytelniającej. Standardowy element architektury zerowego zaufania; w systemach AI obowiązkowy dla dostępu do paneli zarządzania modelami i agentami

SUWERENNOŚĆ I ARCHITEKTURA

Minimalizacja danych (Data Minimisation) *Zasada minimalizacji danych*

Zasada z art. 5 ust. 1 lit. c RODO: dane osobowe muszą być adekwatne, stosowne i ograniczone do tego, co niezbędne do celu przetwarzania. W projektach AI oznacza m.in. ograniczanie zakresu danych treningowych i kontekstu przekazywanego modelom do tego, czego naprawdę wymaga zadanie

REGULACJE I NORMY · W regulacji: RODO

Minimum Sufficient Sovereignty (MSS) *Minimalna wystarczająca suwerenność (MSS)*

Zasada doboru takiego poziomu izolacji, lokalizacji i kontroli nad systemem AI, który wystarcza do spełnienia wymogów regulacyjnych, bezpieczeństwa, ciągłości działania i ochrony własności intelektualnej — i nie jest większy, niż to konieczne. Chroni przed przewymiarowaniem architektury (kosztownym) i jej niedowymiarowaniem (ryzykownym)

SUWERENNOŚĆ I ARCHITEKTURA

Mitigation Response Time

Metryka Cognitive SLA (Tabela 11): czas reakcji na problem kognitywny. Targety: ≤15 min krytyczne, ≤4h standardowe. Metryka Tier 1.

METODYKA CDF

Model Selection Compliance Matrix *Matryca zgodności wyboru modeli AI*

Narzędzie decyzyjne CDF wspierające wybór modeli AI z uwzględnieniem wymogów suwerenności, zgodności regulacyjnej i ograniczeń geopolitycznych. Ocenia modele w wymiarach: licencja, data residency, orbita geopolityczna (US/CN/EU-sovereign), dostępność on-premise, audit trail, SBOM. Powiązany z Dual-Orbit Sovereignty Matrix. Patrz: Załącznik F w Metodocy.

METODYKA CDF

MSS (Minimum Sufficient Sovereignty) *Minimum Sufficient Sovereignty*

Skrót od Minimum Sufficient Sovereignty. Minimalny wymagany poziom suwerenności AI dla danego profilu, określany w F0 na podstawie SLA-S. Patrz: Załącznik F. Zob. Minimum Sufficient Sovereignty (MSS).

METODYKA CDF

Multi-agent *Wieloagentowy system AI*

Architektura, w której wiele wyspecjalizowanych agentów AI współpracuje, aby wykonać złożone zadanie lub cały proces od początku do końca. Zwiększa możliwości, ale mnoży punkty nadzoru — każdy agent potrzebuje tożsamości, uprawnień i logu

PODSTAWY AI

Multi-Jurisdiction Compliance Map *Mapa zgodności wielojurysdykcyjnej*

Artefakt ACE generowany w F0 dla profilu GEO:MULTI: wymogi per jurysdykcja, konflikty regulacyjne, Cross-Border Data Transfer Assessment.

METODYKA CDF

MVOrg / MVOp (dawniej MVO) *MVOrg / MVOp — minimalna organizacja i minimalna operacja*

MVOrg (Minimum Viable Organization) — minimalna konfiguracja kompetencji, nadzoru i zdolności technicznych, jaką organizacja musi mieć przed rozpoczęciem wdrożenia AI: sponsor na poziomie zarządu, osoba odpowiedzialna za wiedzę, inwentaryzacja danych i ocena dojrzałości. MVOp (Minimum Viable Operation) — minimalna, kontrolowana wersja procesu realizowanego przez agentów AI, uruchamiana przed pełnym skalowaniem

WDROŻENIE I ZARZĄDZANIE

MVP *MVP — Minimum Viable Product*

Wersja produktu z minimalnym zestawem funkcji wystarczającym do sprawdzenia pomysłu lub kierunku rozwoju z prawdziwymi użytkownikami. W AI MVP powinien od początku zawierać podstawowe zabezpieczenia i rejestrowanie — dodanie ich później jest kosztowne

WDROŻENIE I ZARZĄDZANIE

N

Named Shortcut (ACE) *Nazwany skrót ACE*

Profil referencyjny ACE (np. STARTER, FINANCE, SOVEREIGN) będący nazwanym skrótem dla typowej kombinacji tagów. Punkt wyjścia konfiguracji – organizacja może modyfikować tagi niezależnie od wybranego profilu.

METODYKA CDF

NASK *NASK — Naukowa i Akademicka Sieć Komputerowa (PIB)*

Państwowy instytut badawczy prowadzący CSIRT NASK — jeden z trzech krajowych zespołów reagowania na incydenty (obok CSIRT GOV i CSIRT MON) — i utrzymujący na zlecenie Ministerstwa Cyfryzacji system S46. Przyjmuje zgłoszenia incydentów od podmiotów

kluczowych i ważnych w Krajowym Systemie Cyberbezpieczeństwa oraz prowadzi badania nad bezpieczeństwem i AI

REGULACJE I NORMY

NASK GOV Prerequisites *Wymagania wstępne NASK GOV*

Checklist wymagań wstępnych dla wariantu CSE-GOV, obejmujący gotowość organizacji do wdrożenia AI w sektorze publicznym RP zgodnie z wytycznymi NASK. Weryfikowany w fazach F0/F1. Obejmuje: zgodność z uKSC/NIS2, dostępność dokumentacji w języku polskim, mapowanie na Strategię Cyberbezpieczeństwa RP 2025–2029, gotowość infrastrukturalną i kadrową. Wprowadzony w CDF 1.4.4 jako element pakietu NASK GOV (Z-090).

METODYKA CDF

NATO PRU *Zasady odpowiedzialnego użycia AI w NATO (Principles of Responsible Use)*

Sześć zasad odpowiedzialnego użycia sztucznej inteligencji przyjętych przez NATO w 2021 r.: zgodność z prawem, odpowiedzialność i rozliczalność, wyjaśnialność i możliwość prześledzenia, niezawodność, sterowność oraz ograniczanie uprzedzeń. Punkt odniesienia dla wdrożeń AI w sektorze obronnym i u dostawców pracujących dla wojska

REGULACJE I NORMY

NGAC *NGAC — Next Generation Access Control*

Model kontroli dostępu oparty na atrybutach i politykach zamiast sztywnych ról (standard ANSI/INCITS 499 i 525; NIST porównał go z XACML w publikacji SP 800-178). Pozwala dynamicznie zarządzać uprawnieniami użytkowników i agentów AI w środowisku wieloagentowym — np. ograniczać dostęp zależnie od kontekstu zadania, klasyfikacji danych i czasu

SUWERENNOŚĆ I ARCHITEKTURA

NIS2/uKSC *Dyrektywa NIS2 i ustawa o Krajowym Systemie Cyberbezpieczeństwa*

Dyrektywa (UE) 2022/2555 (NIS2) i wdrażająca ją ustawa z dnia 23 stycznia 2026 r. o zmianie ustawy o Krajowym Systemie Cyberbezpieczeństwa (Dz.U. 2026 poz. 252), obowiązująca od 3 kwietnia 2026 r. — przepisy o cyberbezpieczeństwie podmiotów kluczowych i ważnych, z okresem dostosowawczym do 3 kwietnia 2027 r. i karami od 3 kwietnia 2028 r. Dla dostawców rozwiązań ICT i AI oznaczają klauzule bezpieczeństwa w umowach, prawo klienta do audytu, zgłaszanie poważnych incydentów (wcześnie ostrzeżenie w 24 godziny, zgłoszenie w 72 godziny) i zarządzanie ryzykiem w łańcuchu dostaw

REGULACJE I NORMY

NIST *NIST — National Institute of Standards and Technology*

Amerykańska agencja federalna przy Departamencie Handlu USA zajmująca się normalizacją i technologią. W obszarze AI i cyberbezpieczeństwa wydaje m.in. NIST AI Risk Management Framework, NIST Cybersecurity Framework, raport AI 800-4 o monitorowaniu wdrożonych systemów AI oraz standardy kryptografii postkwantowej — dokumenty niewiążące, ale powszechnie przyjmowane jako punkt odniesienia

REGULACJE I NORMY

NIST AI 800-4 *NIST AI 800-4 — wyzwania monitorowania wdrożonych systemów AI*

Raport NIST (marzec 2026) analizujący bariery monitorowania systemów AI po wdrożeniu. Porządkuje monitorowanie w sześciu kategoriach: funkcjonalność, eksploatacja, czynnik ludzki, bezpieczeństwo, zgodność i skutki w dużej skali. Nie formułuje wymagań, ale jego taksonomia jest użyteczną podstawą do projektowania nadzoru nad AI w produkcji

REGULACJE I NORMY

Non-repudiation *Niezaprzeczalność (non-repudiation)*

Właściwość systemu gwarantująca, że żaden uczestnik — człowiek ani agent — nie może zaprzeczyć wykonaniu lub zleceniu operacji. Realizowana przez kryptograficznie podpisane wpisy w niemodyfikowalnym dzienniku zdarzeń, opatrzone znacznikami czasu

SUWERENNOŚĆ I ARCHITEKTURA

North Star *North Star — nadrzędny cel programu AI*

Nadrzędny cel strategiczny programu AI: jeden mierzalny wynik biznesowy, nie więcej niż pięć wskaźników KPI, 90-dniowe kamienie milowe i określona koperta inwestycyjna. Zapobiega rozproszonemu portfelowi inicjatyw i służy jako punkt odniesienia przy decyzji o skalowaniu lub zakończeniu pilotażu. W odróżnieniu od „North Star Metric” z metodyk wzrostu produktu (Ellis, 2017) dotyczy transformacji całej organizacji

WDROŻENIE I ZARZĄDZANIE

O

OC *Odpowiedzialność cywilna (ubezpieczenie OC)*

Ubezpieczenie od odpowiedzialności cywilnej. W projektach AI dotyczy polis pokrywających szkody wyrządzone osobom trzecim przez wdrożony system — coraz częściej wymaganych w umowach i zamówieniach publicznych, zwłaszcza po wejściu w życie nowej dyrektywy o odpowiedzialności za produkty

REGULACJE I NORMY

Ochrona informacji niejawnych (OIN) *Ochrona informacji niejawnych*

Wymogi ustawy z dnia 5 sierpnia 2010 r. o ochronie informacji niejawnych: klauzule tajności (zastrzeżone, poufne, tajne, ściśle tajne), poświadczenia bezpieczeństwa personelu, akredytacja systemów teleinformatycznych i warunki przetwarzania. Dla systemów AI w administracji i obronności oznacza to zwykle wdrożenie lokalne, bez połączenia z chmurą publiczną

REGULACJE I NORMY

Odbiór *Odbiór wdrożenia (acceptance)*

Formalne potwierdzenie przez klienta, że wdrożenie spełnia uzgodnione kryteria odbioru — zarówno rekomendowane przez metodykę, jak i własne kryteria klienta. Protokół odbioru zamyka fazę wdrożeniową i jest warunkiem przejścia do eksploatacji

WDROŻENIE I ZARZĄDZANIE

On-premise / On-prem *Wdrożenie lokalne (on-premise)*

Model wdrożenia, w którym infrastruktura AI działa w lokalizacji i pod kontrolą organizacji — we własnej serwerowni lub kolokacji — bez uzależnienia od chmury publicznej jako warstwy podstawowej

SUWERENNOŚĆ I ARCHITEKTURA

Ontology Change Control *Kontrola zmian ontologii*

Rygorystyczna procedura zarządzania zmianami w ontologii grafu wiedzy, obejmująca wersjonowanie i zatwierdzanie zmian. Realizowana przez Knowledge Curator w Fazie 6 CDF, zapewniająca spójność i integralność bazy wiedzy systemu AI.

METODYKA CDF

OpEx *OpEx — wydatki operacyjne*

Bieżące wydatki na utrzymanie i eksploatację systemu AI: energia, licencje i opłaty za tokeny, wsparcie, monitorowanie, ponowne trenowanie modeli. Wraz z wydatkami inwestycyjnymi (CapEx) składa się na całkowity koszt posiadania

WDROŻENIE I ZARZĄDZANIE

Opinia indywidualna *Opinia indywidualna Komisji (ustawa o systemach SI)*

Interpretacja Komisji Rozwoju i Bezpieczeństwa SI dotycząca stosowania AI Act lub ustawy do konkretnego systemu AI, wydawana na wniosek. Wydanie w terminie 30 dni (60 dni w sprawach szczególnie skomplikowanych); brak opinii w terminie oznacza uznanie stanowiska wnioskodawcy. Opinia wiąże Komisję w danej sprawie, a po anonimizacji jest publikowana w BIP

REGULACJE I NORMY · W regulacji: AI Act

Opinia milcząca (SSI) *Milcząca opinia indywidualna*

Skutek niewydania opinii indywidualnej w ustawowym terminie 30 albo 60 dni: uznaje się, że Komisja wydała opinię zgodną ze stanowiskiem przedstawionym we wniosku, czyli korzystną dla wnioskodawcy

REGULACJE I NORMY

Orchestration *Orkiestracja*

Koordynowanie przepływu zadań między wieloma agentami AI, systemami lub komponentami w celu realizacji złożonego procesu. Warstwa orkiestracji decyduje, kto co robi i w jakiej kolejności — dlatego każda jej decyzja powinna trafiać do niezmiennego dziennika zdarzeń

PODSTAWY AI

Organ notyfikujący (SSI) *Organ notyfikujący (AI Act, ustawa o systemach SI)*

Organ odpowiedzialny za notyfikowanie jednostek oceniających zgodność systemów AI i nadzór nad nimi (art. 28 AI Act). W Polsce jest nim minister właściwy do spraw informatyzacji, a Polskie Centrum Akredytacji współtworzy program akredytacji i informuje ministra o istotnych decyzjach akredytacyjnych

REGULACJE I NORMY · W regulacji: AI Act

OWASP *OWASP — Open Worldwide Application Security Project*

Międzynarodowa organizacja non-profit publikująca otwarte standardy i listy zagrożeń bezpieczeństwa aplikacji. W obszarze AI wydaje m.in. OWASP GenAI LLM Top 10 (edycja 2026) oraz OWASP Top 10 for Agentic Applications — powszechnie stosowane punkty odniesienia przy testach bezpieczeństwa systemów opartych na modelach językowych i agentach

REGULACJE I NORMY

OWASP ASI Mapping Matrix *OWASP Agentic Security Initiative Mapping Matrix*

Tabela mapowania 10 ryzyk OWASP Top 10 for Agentic Applications (ASI01–ASI10, edycja 2026) na kontrole, artefakty i fazy CDF. Obejmuje: ASI01 Agent Goal Hijack, ASI02 Tool Misuse & Exploitation, ASI03 Identity & Privilege Abuse, ASI04 Agentic Supply Chain Vulnerabilities, ASI05 Unexpected Code Execution, ASI06 Memory & Context Poisoning, ASI07 Insecure Inter-Agent Communication, ASI08 Cascading Failures, ASI09 Human-Agent Trust Exploitation, ASI10 Rogue Agents. Patrz: Załącznik B.

METODYKA CDF

P

PAG *Public Administration Guardrails*

Rodzina trzech kontroli CDF dla administracji publicznej, wprowadzona w 1.6.0 i bramkowana tagiem GOV; moduł krajowy aktywuje tag GEO:PL. Obejmuje PAG-1 Necessity Gate, PAG-2 Human-Contact Fallback i PAG-3 Approved Tooling Register. Wszystkie kontrole mają charakter blokujący: niespełnienie zatrzymuje przejście bramki fazowej. PAG jest rodziną kontroli, nie tagiem konfiguracyjnym.

METODYKA CDF

PAG-1 *Necessity Gate – bramka zasadności*

Kontrola CDF (KD1, KD8; faza F0) rodziny PAG. Bramka zasadności zastosowania AI oparta na Etapie 0 Przewodnika MC po AI. Wyjście trójwartościowe: GO – uzasadnienie przyjęte; NO-GO – cel osiągalny prostszymi środkami, projekt zamknięty w F0; REDESIGN – zakres wymaga przeprojektowania. Wariant NO-GO jest równoprawnym wynikiem fazy F0. Artefakt A-GOV-01. W profilach innych niż GOV kontrola jest rekomendowana, nie obligatoryjna.

METODYKA CDF

PAG-2 *Human-Contact Fallback – zastępcza ścieżka kontaktu z człowiekiem*

Kontrola CDF (KD9, KD7; fazy F2 i F4) rodziny PAG. Każdy proces styku AI z obywatelem musi udostępniać alternatywny kanał obsługi prowadzący do pracownika: funkcję przełączenia do człowieka w każdym punkcie interakcji, kanał niezależny od systemu AI utrzymywany w deklarowanych godzinach oraz ścieżkę zgłoszenia zastrzeżeń, odrębną od kanału zgłaszania obaw AI z F6. Brak zastępczej ścieżki blokuje bramkę wyjścia z F4. Artefakt A-F2-16.

METODYKA CDF

PAG-3 *Approved Tooling Register – rejestr narzędzi zatwierdzonych*

Kontrola CDF (KD6, KD1; fazy F2 i F6) rodziny PAG. Rejestr narzędzi AI dopuszczonych do użytku służbowego wraz z polityką korzystania z AI. Różni się od Shadow Agent Governance: SAG-4 wykrywa nieautoryzowanych agentów w architekturze systemu, PAG-3 obejmuje narzędzia GenAI używane przez pracowników poza architekturą wdrożenia. Obie kontrole zasilają wspólny przegląd w F6. Artefakty A-GOV-02 i A-GOV-03.

METODYKA CDF

PCOC (Połączone Centrum Operacyjne Cyberbezpieczeństwa) *Połączone Centrum Operacyjne Cyberbezpieczeństwa (PCOC)*

Połączone Centrum Operacyjne Cyberbezpieczeństwa — przewidziane w Strategii Cyberbezpieczeństwa RP na lata 2025–2029 (uchwała Rady Ministrów z 10 marca 2026 r.) jako załączek centralnej instytucji koordynującej cyberbezpieczeństwo na poziomie krajowym: wspólne miejsce przyjmowania zgłoszeń i koordynacji reakcji na incydenty między zespołami CSIRT

REGULACJE I NORMY

Phase Gate *Bramka fazowa (phase gate)*

Formalny punkt kontrolny na przejściu między fazami projektu. Określa kryteria wejścia, wymagane artefakty i kryteria wyjścia (definicję ukończenia); przejście wymaga zatwierdzenia przez osobę odpowiedzialną z macierzy RACI. W metodyce CDF bramki oddzielają każdą z faz — od rozpoznania po eksploatację

WDROŻENIE I ZARZĄDZANIE

Phase Rollback *Nawrót fazowy*

Mechanizm powrotu do poprzedniej fazy CDF w przypadku niespełnienia kryteriów wyjścia (DoD) bieżącej fazy. Phase Rollback wymaga formalnej decyzji osoby Accountable i dokumentacji przyczyn nawrotu. Źródło: Metodyka, sekcja Kryteria wejścia i wyjścia faz CDF. Wprowadzone w v1.4.4.

METODYKA CDF

Piaskownica regulacyjna AI

Kontrolowane środowisko, w którym innowacyjne systemy AI można rozwijać i testować pod nadzorem organu przed wprowadzeniem na rynek (art. 57 AI Act). W Polsce piaskownicę prowadzi Komisja Rozwoju i Bezpieczeństwa SI: udział trwa 6-12 miesięcy, uczestników wybiera się w konkursie, a dla MSP udział jest nieodpłatny. Uczestnik składa raport i może wystąpić o opinię indywidualną

REGULACJE I NORMY · W regulacji: AI Act

PKG *Personal Knowledge Graph*

Osobisty graf wiedzy pracownika – ustrukturyzowana reprezentacja relacji między pojęciami, dokumentami, procesami i kontekstami, którymi operuje dany pracownik. Komponent architektury Cyfrowego Bliźniaka Poznawczego (CBP) obok Personal Language Model (PLM) i Profilu Interpretacyjnego (PI).

METODYKA CDF

Plan monitoringu Technologies-in-Practice *Plan monitorowania praktyk*

technologicznych

Artefakt Fazy 3 służący obserwacji tego, jak pracownicy realnie korzystają z AI w praktyce oraz czy nie utrwalają się nieoptymalne, niebezpieczne lub niezgodne wzorce użycia.

METODYKA CDF

PLD (dyrektywa 2024/2853) *Dyrektywa o odpowiedzialności za produkty wadliwe — (UE) 2024/2853*

Dyrektywa (UE) 2024/2853 w sprawie odpowiedzialności za produkty wadliwe, zastępująca dyrektywę z 1985 r. Po raz pierwszy wprost obejmuje oprogramowanie i systemy AI jako produkt, wprowadza odpowiedzialność niezależną od winy, obowiązek ujawnienia dowodów (art. 9) oraz domniemania wadliwości i związku przyczynowego (art. 10). Państwa członkowskie muszą ją wdrożyć do 9 grudnia 2026 r.

REGULACJE I NORMY · W regulacji: PLD 2024/2853

PLM *Personal Language Model*

Skrót od Personal Language Model. Zob. PLM (Personal Language Model).

METODYKA CDF

PLM (Personal Language Model) *Osobisty model językowy*

Indywidualnie dostrojony model językowy — składnik Cyfrowego Bliźniaka Poznawczego — kalibrowany na stylu komunikacji, terminologii i sposobie rozumowania konkretnego pracownika. Kalibracja obejmuje wyłącznie dane związane z rolą zawodową; dane prywatne i szczególnych kategorii (art. 9 RODO) są wyłączone, a przetwarzanie wymaga oceny skutków (DPIA, art. 35 RODO)

PODSTAWY AI · W regulacji: RODO

Post-Deployment Monitoring *Monitorowanie po wdrożeniu*

Ciągłe monitorowanie systemu AI po uruchomieniu produkcyjnym w sześciu kategoriach opisanych przez NIST AI 800-4: funkcjonalność, eksploatacja, czynnik ludzki, bezpieczeństwo, zgodność regulacyjna i skutki w dużej skali. Dla systemów wysokiego ryzyka jest to obowiązek dostawcy z art. 72 AI Act

WDROŻENIE I ZARZĄDZANIE · W regulacji: AI Act

Prawo do usunięcia (Right to Erasure) *Prawo do usunięcia danych („prawo do bycia zapomnianym”)*

Prawo osoby, której dane dotyczą, do żądania usunięcia jej danych osobowych (art. 17 RODO). W systemach AI wymaga przemyślenia zawczasu: dane trzeba móc usunąć z baz wektorowych, pamięci podręcznej, logów i zbiorów do strojenia modelu, a nie tylko z bazy źródłowej

REGULACJE I NORMY · W regulacji: RODO

PRODUCTION (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla wdrożeń produkcyjnych lub transformacji. Włącza pełne fazy F0-F4 (lub F0-F5 przy SCALE) oraz F6-Lite. NIE aktywuje wymogów compliance – tylko pełny flow wdrożeniowy.

Production Cost Model *Model kosztów produkcyjnych*

Oszacowanie pełnych kosztów przejścia z pilotażu AI do środowiska produkcyjnego: infrastruktury, integracji, zgodności regulacyjnej, obsady operacyjnej, budowy bazy wiedzy i zabezpieczeń. Doświadczenie pokazuje, że produkcja kosztuje wielokrotność pilotażu — dlatego model liczy się przed decyzją o skalowaniu, a nie po niej

WDROŻENIE I ZARZĄDZANIE

Profil HEALTH *Profil referencyjny HEALTH*

Profil referencyjny ACE dla podmiotów ochrony zdrowia i dostawców AI medycznego. Archetyp: dane zdrowotne, wyrób medyczny AI, EHDS. Tagi aktywne: CORE + PRODUCTION + REGULATED + HEALTH + GEO:xx (+ AGENT gdy agenty). Fazy F0→F6-Lite, Cognitive SLA Tier 2, HCG-Full. Jedenasty profil referencyjny CDF. Karta: Tabela 45a w Metodyce. Wprowadzony w CDF 1.5.0.

METODYKA CDF

Prompt *Polecenie (zapytanie) do AI*

Tekst przekazywany do modelu językowego lub agenta AI, żeby wykonał zadanie, udzielił odpowiedzi albo uruchomił procedurę. W systemach firmowych prompt to także element konfiguracji — jego treść podlega wersjonowaniu i kontroli jak kod

PODSTAWY AI

Prompt Injection *Wstrzyknięcie polecenia*

Atak, w którym do danych przetwarzanych przez model — dokumentu, e-maila, strony WWW, wyniku narzędzia — wstrzykiwane są ukryte instrukcje, tak by model zrobił coś, czego użytkownik nie zlecił: ujawnił dane, pominął zasady, wykonał akcję. Pozycja LLM01 na liście OWASP GenAI LLM Top 10 i punkt wyjścia dla ASI01 (przejęcie celu agenta) na liście OWASP dla aplikacji agentowych

PODSTAWY AI

Prompt Origin

Metadane źródła promptu: direct (użytkownik), template, agent, trigger. Element Cognitive Observability w F6.

METODYKA CDF

Proporcjonalność (Proportionality) *Zasada proporcjonalności (DORA)*

Zasada z art. 4 DORA: wymogi stosuje się proporcjonalnie do wielkości, profilu ryzyka oraz charakteru, skali i złożoności działalności podmiotu finansowego. Mikroprzedsiębiorstwa korzystają z licznych ułatwień, a wskazane kategorie małych podmiotów — z uproszczonych ram zarządzania ryzykiem ICT (art. 16)

REGULACJE I NORMY · W regulacji: DORA

Przewodnik MC po AI *Przewodnik po sztucznej inteligencji dla administracji publicznej (MC)*

Dokument Ministerstwa Cyfryzacji opublikowany 5 marca 2026 r. dla jednostek administracji publicznej wdrażających AI. Nie ma mocy wiążącej, ale wyznacza standard postępowania, do którego odwołują się organy nadzoru i audytorzy: proces wdrożenia w ośmiu etapach — od oceny zasadności (Etap 0) po przegląd okresowy (Etap 7) — zestawienie „co robić / czego nie robić” w sześciu obszarach oraz sześcioletnia checklista dla pracownika

REGULACJE I NORMY

Psychological Readiness *Gotowość psychologiczna organizacji*

Wymiar dojrzałości organizacji oceniany przed wdrożeniem AI: nastawienie pracowników, akceptacja zmiany, odporność na obciążenie poznawcze i gotowość do współpracy z narzędziami AI. Niska gotowość psychologiczna przewiduje niską adopcję niezależnie od jakości technologii

WDROŻENIE I ZARZĄDZANIE

Psychological Safety Protocol *Protokół bezpieczeństwa psychologicznego*

Zestaw działań wdrażanych w Fazie 3, obejmujący pomiar bezpieczeństwa psychologicznego zespołu, przeciwdziałanie stygmatyzacji użytkowników AI, interwencje zarządzania zmianą oraz monitoring utrwalania nieoptymalnych praktyk użycia AI.

METODYKA CDF

Q

Quality Management System (System zarządzania jakością) *System*

zarządzania jakością AI (art. 17 AI Act)

System wymagany przez art. 17 AI Act od dostawców systemów AI wysokiego ryzyka. Obejmuje strategię zgodności regulacyjnej, procedury projektowania, rozwoju, testowania i walidacji, zarządzanie danymi i ryzykiem, monitorowanie po wprowadzeniu do obrotu, zgłaszanie incydentów i rozliczalność kierownictwa. Może być zintegrowany z istniejącym systemem ISO 9001 lub ISO/IEC 42001

REGULACJE I NORMY · W regulacji: AI Act · ISO/IEC 42001

R

RACI *Macierz odpowiedzialności RACI*

Macierz przypisująca każdemu zadaniu rolę: R (Responsible — wykonuje), A (Accountable — odpowiada i zatwierdza), C (Consulted — konsultowany), I (Informed — informowany). We wdrożeniach AI obejmuje zwykle sponsora, osobę odpowiedzialną za zgodność, lidera wdrożenia, architekta, bezpieczeństwo i właściciela procesu biznesowego

WDROŻENIE I ZARZĄDZANIE

Rada przy Komisji (SSI) *Rada — organ opiniotawczo-doradczy Komisji (ustawa o systemach SI)*

Organ opiniotawczo-doradczy przy Komisji Rozwoju i Bezpieczeństwa SI, w ustawie nazwany po prostu Radą, liczący od 9 do 15 członków wybieranych przez Komisję na dwuletnią kadencję.

Kandydatów zgłaszają m.in. strona samorządowa Komisji Wspólnej Rządu i Samorządu Terytorialnego, Rzecznik Praw Obywatelskich, Rzecznik Praw Dziecka, Rzecznik Praw Pacjenta, Rzecznik MŚP i Główny Inspektor Pracy. Opinie Rady nie wiążą Komisji

REGULACJE I NORMY

RAG *Retrieval-Augmented Generation — generowanie wspomagane wyszukiwaniem*

Architektura, w której model generatywny przed odpowiedzią wyszukuje fragmenty w kontrolowanej bazie wiedzy i na nich opiera odpowiedź. Dzięki temu odpowiedzi są aktualne, oparte na źródłach organizacji i możliwe do zweryfikowania, a ryzyko halucynacji spada

PODSTAWY AI

Raport bazowy bezpieczeństwa psychologicznego

Artefakt Fazy 3 dokumentujący wyjściowy poziom bezpieczeństwa psychologicznego zespołów przed wdrożeniem szerokiej współpracy człowiek-AI.

METODYKA CDF

RBAC *RBAC — kontrola dostępu oparta na rolach*

Model kontroli dostępu, w którym uprawnienia przypisuje się rolom organizacyjnym, a nie pojedynczym osobom. Podstawowy mechanizm zarządzania uprawnieniami użytkowników i agentów AI; w środowiskach wieloagentowych uzupełniany o polityki atrybutowe (ABAC, NGAC)

SUWERENNOŚĆ I ARCHITEKTURA

Reasoning Accuracy Rate *Wskaźnik trafności rozumowania*

Odsetek odpowiedzi lub decyzji systemu AI zgodnych z prawdą referencyjną albo z oczekiwanym stanem wiedzy w danej dziedzinie. Podstawowa metryka jakości rozumowania — mierzona na zestawie testowym, nie na wrażeniu użytkowników

PODSTAWY AI

Recovery Rate *Wskaźnik odzyskiwania*

Metryka warstwy Agent SLA (Tabela 52 CDF) mierząca zdolność agenta do samodzielnego przywrócenia poprawnego działania po wykryciu błędu lub incydentu kognitywnego, bez interwencji człowieka.

METODYKA CDF

REDUCTION (flag) *Flaga redukcji zakresu*

Flaga w ACE Profile Version Record sygnalizująca, że zmiana profilu była Scope Reduction (obniżenie zakresu). Wymaga pisemnego uzasadnienia i weryfikacji regulatory floor.

METODYKA CDF

REGULATED (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla sektora regulowanego lub gdy wdrożenie przetwarza dane osobowe/operacyjne. Włącza sekcje compliance, audytu, compliance-first delivery (Faza 1.5). Aktywowane przez sektor, dane lub lokalizację.

METODYKA CDF

Rejestr czynności przetwarzania (ROPA) *Rejestr czynności przetwarzania (RCP, ang. ROPA)*

Rejestr prowadzony przez administratora danych zgodnie z art. 30 RODO, dokumentujący cele przetwarzania, kategorie osób i danych, odbiorców, przekazywanie do państw trzecich, planowane terminy usunięcia i środki bezpieczeństwa. Każdy system AI przetwarzający dane osobowe — w tym asystent korzystający z dokumentów firmowych — powinien mieć w nim swój wpis

REGULACJE I NORMY · W regulacji: RODO

Rejestr zasobów wielokrotnego użytku *Rejestr zasobów wielokrotnego użycia*

Zbiór komponentów, wzorców, agentów, szablonów i artefaktów, które mogą być bezpiecznie reużyte w kolejnych wdrożeniach AI w ramach organizacji. Element Fazy 5 CDF.

METODYKA CDF

Release Quality Gate *Checklista jakości wydania*

11-punktowa checklista kontroli jakości, którą każda wersja CDF musi przejść przed publikacją. Obejmuje: referencje, numerację, spójność faz, artefaktów, słownika, załączników, dat, claimów, skal, formatowania i triady wydania. Owner: Delivery Lead. Zatwierdzenie: Sponsor. Źródło: Metodyka, sekcja Checklista jakości wydania CDF. Wprowadzone w v1.4.4.

METODYKA CDF

Retraining *Ponowne trenowanie (dotrenowanie) modelu*

Aktualizacja modelu AI na nowych danych w celu poprawy skuteczności, dostosowania do zmian w dziedzinie lub ograniczenia dryfu. W systemach regulowanych każde dotrenowanie może być znaczącą modyfikacją — z konsekwencjami dla odpowiedzialności i dokumentacji

PODSTAWY AI

Roadmap Dojrzałości *Mapa drogowa dojrzałości AI*

Ścieżka rozwoju organizacji od obecnego poziomu dojrzałości AI do wyższego, z harmonogramem i kamieniami milowymi — np. od pierwszych pilotaży do standardowego, nadzorowanego wykorzystania AI w ciągu sześciu miesięcy

WDROŻENIE I ZARZĄDZANIE

RODO *RODO — ogólne rozporządzenie o ochronie danych (UE) 2016/679*

Polska nazwa ogólnego rozporządzenia o ochronie danych (UE) 2016/679, stosowanego bezpośrednio od 25 maja 2018 r. Dla systemów AI kluczowe są zasady zgodności z prawem, minimalizacji i ograniczenia celu, ocena skutków (art. 35), prawa osób oraz art. 22 ograniczający decyzje oparte wyłącznie na zautomatyzowanym przetwarzaniu, w tym profilowaniu. Zob. GDPR

REGULACJE I NORMY

ROI *ROI — zwrot z inwestycji*

Wskaźnik finansowy mierzący opłacalność projektu przez zestawienie uzyskanych korzyści z poniesionymi kosztami. W projektach AI korzyści warto liczyć ostrożnie i z uwzględnieniem przejściowego spadku produktywności na starcie (krzywa J)

WDROŻENIE I ZARZĄDZANIE

ROI Baseline Report *Raport bazowy ROI*

Artefakt Fazy 0 określający stan wyjściowy dla pomiaru korzyści biznesowych, uwzględniający również wpływ krzywej J produktywności.

METODYKA CDF

Root Cause Analysis *Analiza przyczyn źródłowych*

Systematyczne badanie przyczyn incydentu lub błędu systemu AI po to, by usunąć źródło problemu, a nie tylko jego objaw. W praktyce AI odpowiada na pytanie, czy zawiódł model, dane, prompt, integracja czy nadzór

PODSTAWY AI

Rozporządzenie (UE) 2026/1744 *Rozporządzenie (UE) 2026/1744 — Digital Omnibus on AI*

Rozporządzenie nowelizujące AI Act (2024/1689) oraz rozporządzenia 2018/1139 i 2023/1230, ogłoszone w Dz.Urz. UE 24.07.2026, w mocy od 27.07.2026. Źródło aktualnych terminów stosowania AI Act. Zob. Digital Omnibus (EU)

REGULACJE I NORMY · W regulacji: AI Act

Rozporządzenie maszynowe (rozp. 2023/1230) *Rozporządzenie w sprawie maszyn — (UE) 2023/1230*

Rozporządzenie (UE) 2023/1230 w sprawie maszyn, zastępujące dyrektywę 2006/42/WE, stosowane od 20 stycznia 2027 r. Obejmuje maszyny, w których funkcje bezpieczeństwa realizuje oprogramowanie z elementami AI; takie systemy są jednocześnie systemami wysokiego ryzyka w rozumieniu załącznika I AI Act, z obowiązkami stosowanymi od 02.08.2028

REGULACJE I NORMY · W regulacji: AI Act

S

SA-XAI *SA-XAI — wyjaśnialność oparta na świadomości sytuacyjnej*

Podejście do wyjaśnialności AI, które mapuje decyzje systemu na trzy poziomy świadomości sytuacyjnej operatora: percepcję (co system widzi), rozumienie (jak to interpretuje) i projekcję (co z tego wynika). Daje osobie nadzorującej kontekst potrzebny do świadomego zatwierdzenia lub odrzucenia decyzji

WDROŻENIE I ZARZĄDZANIE

SAG (Shadow Agent Governance) *Kontrola zarządzania agentami nieautoryzowanymi*

Akronim oznaczający zestaw pięciu kontroli CDF adresujących ryzyko nieautoryzowanych agentów AI: SAG-1 Discovery Scan, SAG-2 Risk Triage, SAG-3 Remediation Path, SAG-4 Continuous Monitoring, SAG-5 Reporting. Zob. Shadow Agent Governance.

METODYKA CDF

SAVANT *SAVANT (nazwa własna produktu)*

SAVANT to nazwa własna produktu allclouds.pl – suwerenny system kognitywny klasy enterprise. Nie wymaga definicji akronimowej. www.savant-ai.app

METODYKA CDF

SBOM (Software Bill of Materials) *SBOM — spis komponentów oprogramowania*

Formalna inwentaryzacja komponentów, z których zbudowany jest produkt cyfrowy: bibliotek, frameworków, interfejsów API i zależności. Wymagana przez Cyber Resilience Act jako podstawa monitorowania podatności; w praktyce rozszerzana o modele AI, modele osadzeń i bazy wektorowe (AI BOM), co pozwala szybko reagować na luki w łańcuchu dostaw

REGULACJE I NORMY · W regulacji: CRA

SCALE (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla organizacji enterprise (5000+ pracowników), transformacji lub wdrożeń multi-agent. Włącza sekcje skalowania, industrializacji, CBP, CogOps-Full, F6-Full.

METODYKA CDF

Scale Path Definition *Definicja ścieżki skalowania*

Obowiązkowy artefakt Fazy 0 CDF precyzujący exit criteria pilotażu, ścieżkę do wdrożenia produkcyjnego, Production Cost Model oraz kryteria kill/pivot. Ma zapobiegać utknięciu organizacji na etapie pilotaży.

METODYKA CDF

Scale-or-Kill Gate *Brama decyzyjna Scale-or-Kill*

Obowiązkowy, binarny punkt decyzyjny na końcu pilotażu AI, który celowo wyklucza opcję „przedłużmy pilotaż”. Dwa wyjścia: skalowanie do produkcji albo zakończenie. Zakończenie jest wynikiem pozytywnym — organizacja unika kosztownej inwestycji produkcyjnej w rozwiązanie, które nie dowiodło wartości

WDROŻENIE I ZARZĄDZANIE

SCIM *SCIM — System for Cross-domain Identity Management*

Otwarty standard (RFC 7643 i 7644) automatyzujący zakładanie, aktualizację i usuwanie tożsamości między systemami. W środowiskach AI stosowany do spójnego nadawania i odbierania uprawnień użytkownikom oraz agentom — tak, aby odejście pracownika lub wycofanie agenta natychmiast zamykało dostęp

SUWERENNOŚĆ I ARCHITEKTURA

Scope Reduction *Redukcja zakresu wdrożenia*

Świadome ograniczenie ambicji wdrożenia AI — np. rezygnacja z transformacji całego procesu na rzecz pilotażu. Dopuszczalne pod trzema warunkami: pisemne uzasadnienie, zachowanie minimum regulacyjnego (wymogi obowiązkowe nie mogą zostać wyłączone) i odnotowanie zmiany w historii wersji konfiguracji

WDROŻENIE I ZARZĄDZANIE

Scoring Gotowości AI *Ocena gotowości organizacji do AI (AI readiness scoring)*

Ilościowa ocena dojrzałości organizacji do wdrożenia AI w wymiarach: dane, nadzór i zgodność, kompetencje oraz gotowość psychologiczna. Wynik wskazuje, od czego zacząć i jakie luki zamknąć przed pilotażem

WDROŻENIE I ZARZĄDZANIE

Security by Default *Bezpieczeństwo domyślne (security by default)*

Zasada projektowa, zgodnie z którą bezpieczne ustawienia, polityki i ograniczenia są aktywne od razu po instalacji, a nie dopiero po ręcznej konfiguracji. Dla systemów AI oznacza m.in. minimalne uprawnienia narzędzi agenta, wyłączone domyślnie połączenia zewnętrzne i włączone rejestrowanie

SUWERENNOŚĆ I ARCHITEKTURA

Security by Design *Bezpieczeństwo w fazie projektowania (security by design)*

Zasada wbudowywania środków bezpieczeństwa w system od początku projektu — na etapie architektury i wymagań — a nie dodawania ich po wdrożeniu. Wymóg Cyber Resilience Act wobec produktów z elementami cyfrowymi; w AI obejmuje m.in. modelowanie zagrożeń dla agentów i ochronę danych treningowych

SUWERENNOŚĆ I ARCHITEKTURA · W regulacji: CRA

Shadow Agent Governance *Zarządzanie nieautoryzowanymi agentami AI*

Zestaw praktyk wykrywania i porządkowania agentów AI, które działają w organizacji poza formalnym nadzorem: inwentaryzacja, ocena ryzyka, legalizacja albo wyłączenie, ciągły monitoring i raportowanie. Rozszerzenie problemu Shadow AI na środowisko wieloagentowe

PODSTAWY AI

Shadow AI *Nieautoryzowane użycie AI*

Korzystanie przez pracowników z publicznych narzędzi AI poza wiedzą i kontrolą organizacji — wklejanie do czatbota umów, danych osobowych, kodu. Najczęstsza droga wycieku danych związana z AI i pierwszy problem, który rozwiązuje polityka korzystania z AI wraz z rejestrem zatwierdzonych narzędzi

PODSTAWY AI

SIEM *SIEM — Security Information and Event Management*

System zarządzania informacjami i zdarzeniami bezpieczeństwa, który agreguje logi i alerty z wielu źródeł, koreluje je i wykrywa zagrożenia oraz incydenty. Logi systemów AI i agentów powinny do niego trafiać na równi z logami sieci i aplikacji

SUWERENNOŚĆ I ARCHITEKTURA

Single Agent *Pojedynczy agent AI*

Agent AI odpowiedzialny za jedno wyodrębnione zadanie, na przykład klasyfikację dokumentów, ekstrakcję danych z faktur lub przygotowanie rekomendacji. Najprostsza i najłatwiejsza do nadzoru forma wdrożenia agentowego

PODSTAWY AI

Skala Autonomii (L0-L4) *Pięciopoziomowa skala autonomii agentów AI*

Skala określająca zakres uprawnień decyzyjnych agenta: L0 — agent tylko doradza, wszystkie decyzje i działania podejmuje człowiek; L1 (human-in-the-loop) — agent wykonuje, człowiek zatwierdza każde działanie; L2 (human-on-the-loop) — agent wykonuje, człowiek nadzoruje i może przerwać; L3 (human-over-the-loop) — agent działa samodzielnie w zdefiniowanych ramach; L4 — pełna autonomia, wyłącznie w niekrytycznych procesach zaplecza

PODSTAWY AI

Skale i taksonomie CDF *CDF Scales and Taxonomies*

Blok definicyjny wprowadzony w 1.6.0, jedyne źródło poziomów skal CDF: LOA (L0-L4), CFL (CFL-1 do CFL-5), guardrails ACE (BLOCKING, AUTO-CORRECT, ADVISORY), tiery Cognitive SLA, tiery HCG, tiery F6 CogOps, skala degradacji kognitywnej, eskalacja Cognitive SLA, poziomy suwerenności AI (Poziom 1-3) i poziomy suwerenności compute (S1-S4). Pozostałe sekcje metodyki odwołują się do poziomu i nie powtarzają definicji ani liczby elementów.

METODYKA CDF

Skill Partnership *Partnerstwo kompetencyjne człowiek-AI*

Model współpracy, w którym człowiek i agent AI uzupełniają swoje kompetencje: AI przejmuje zadania powtarzalne i analityczne, człowiek wnosi osąd, kontekst i odpowiedzialność. Opisywany w postaci matrycy „kto co robi” dla każdego procesu

WDROŻENIE I ZARZĄDZANIE

SLA *SLA — umowa o poziomie usług*

Umowa o poziomie usług określająca gwarantowane parametry działania — dostępność, czas reakcji, czas naprawy, jakość obsługi — oraz konsekwencje ich niedotrzymania. Dla usług AI warto rozszerzyć ją o parametry jakości odpowiedzi, np. dopuszczalny odsetek błędów

SUWERENNOŚĆ I ARCHITEKTURA

SLA-S *Sovereignty Level Assessment*

Skrót używany dla oceny poziomu suwerenności realizowanej w Fazie 0, służącej do przypisania właściwego wzorca wdrożeniowego AI dla danego obszaru biznesowego.

METODYKA CDF

SoA *Deklaracja stosowania (Statement of Applicability, SoA)*

Dokument wymagany przy wdrożeniu i certyfikacji systemu zarządzania (np. ISO/IEC 42001 lub ISO/IEC 27001), określający, które kontrole z załącznika A normy organizacja stosuje, a które wyłącza — wraz z uzasadnieniem i sposobem wdrożenia. Punkt wyjścia dla audytora

REGULACJE I NORMY · W regulacji: ISO/IEC 42001

SOKiK (właściwość AI) *Sąd Ochrony Konkurencji i Konsumentów — sprawy z zakresu AI*

Sąd Okręgowy w Warszawie — Sąd Ochrony Konkurencji i Konsumentów, właściwy do rozpoznawania odwołań od decyzji i zażaleń na postanowienia Komisji Rozwoju i Bezpieczeństwa SI, w odrębnym postępowaniu w sprawach z zakresu sztucznej inteligencji wprowadzonym do Kodeksu postępowania cywilnego ustawą o systemach SI

REGULACJE I NORMY

SOVEREIGN (tag ACE) *Tag konfiguracyjny ACE*

Tag ACE aktywowany dla wdrożeń z informacjami niejawnymi lub infrastrukturą krytyczną. Włącza sekcję pełnej suwerenności: air-gap, EuroHPC, Hardware Binding. Automatycznie wymusza REGULATED + GOV.

METODYKA CDF

Sovereign AI *Suwerenna AI (sovereign AI)*

Model wdrożenia i eksploatacji sztucznej inteligencji, w którym organizacja zachowuje wysoką lub pełną kontrolę nad danymi, modelem, logami, integracjami i warstwą operacyjną — niezależnie od tego, czy infrastruktura działa lokalnie, czy w zaufanej chmurze w jurysdykcji UE

SUWERENNOŚĆ I ARCHITEKTURA

Sovereign AI Rationale *Uzasadnienie suwerennej architektury AI*

Artefakt fazy F0 dokumentujący strategiczne, regulacyjne i operacyjne przesłanki wyboru suwerennej architektury AI dla danego wdrożenia. Obejmuje: analizę ryzyk geopolitycznych, wymogi data residency, zgodność z EU AI Act i NIS2/uKSC, ocenę Dual-Orbit Sovereignty Matrix oraz uzasadnienie poziomu suwerenności (MSS). Wymagany dla profili GOV, SOVEREIGN i DEFENSE. Wprowadzony w CDF 1.4.4 (Z-089/Z-090).

METODYKA CDF

Sovereign Compute *Suwerenne zasoby obliczeniowe (sovereign compute)*

Infrastruktura obliczeniowa pozostająca pod pełną kontrolą prawną i operacyjną organizacji lub państwa — od własnej serwerowni, przez chmurę w jurysdykcji UE, po układy hybrydowe. Wymagany poziom kontroli wynika z klasyfikacji danych, ekspozycji regulacyjnej i krytyczności systemu

SUWERENNOŚĆ I ARCHITEKTURA

Sovereign Deployment Pattern *Wzorzec wdrożenia suwerennego*

Ostateczny wybór modelu operacyjnego infrastruktury AI dokonywany w Fazie 1 CDF na podstawie Data Gravity Assessment i Sovereignty Level Assessment, spośród trzech opcji: Fully On-Premise (pełna suwerenność), Hybrid Sovereign-Cloud (dane krytyczne lokalnie, obliczenia w chmurze) i Air-Gapped Defense (fizyczna izolacja sieciowa dla sektora obronnego).

METODYKA CDF

Sovereignty Dimensions *Wymiary suwerenności AI*

Cztery kryteria oceny, jakiego poziomu suwerenności wymaga system AI: klasyfikacja danych, ekspozycja regulacyjna, krytyczność operacyjna oraz zależność od łańcucha dostaw. Ich łączna

ocena wskazuje, czy wystarczy chmura publiczna, czy potrzebne jest wdrożenie lokalne albo izolowane

SUWERENNOŚĆ I ARCHITEKTURA

Sovereignty Level Assessment (SLA-S) *Ocena poziomu suwerenności*

Narzędzie Fazy 0 badające dany obszar biznesowy w czterech wymiarach w celu przypisania go do odpowiedniego poziomu i wzorca wdrożeniowego AI. Chroni organizację przed przewymiarowaniem lub niedoszacowaniem poziomu izolacji i kontroli.

METODYKA CDF

SPIFFE/SPIRE *SPIFFE / SPIRE — tożsamość obciążeń*

Otwarty standard CNCF (Secure Production Identity Framework for Everyone) i jego implementacja referencyjna SPIRE do nadawania kryptograficznie weryfikowalnych tożsamości usługom i procesom w środowiskach rozproszonych. Pozwala identyfikować agentów AI bez współdzielonych sekretów i haseł

SUWERENNOŚĆ I ARCHITEKTURA

Statement of Applicability *Statement of Applicability (SoA)*

Zob. SoA — deklaracja stosowania

REGULACJE I NORMY

Status standardu *Standard Status*

Kolumna w Matrycy Compliance (Załącznik E.1) określająca bieżący status danej regulacji lub standardu. Możliwe wartości: Published (opublikowany i obowiązujący), Draft (projekt w konsultacjach), Pending (oczekuje na zatwierdzenie lub ogłoszenie), Reserved (mapowanie zarezerwowane do czasu publikacji standardu). Wprowadzona w ramach Z-010 dla lepszej orientacji w dojrzałości regulacyjnej.

METODYKA CDF

Strategia Cyberbezpieczeństwa RP 2025-2029 *Strategia Cyberbezpieczeństwa Rzeczypospolitej Polskiej na lata 2025-2029*

Dokument strategiczny przyjęty uchwałą Rady Ministrów z 10 marca 2026 r. (M.P. poz. 309), wyznaczający sześć celów szczegółowych wzmacniających Krajowy System Cyberbezpieczeństwa — m.in. odporność administracji i sektorów kluczowych, suwerenność technologiczną, kompetencje kadr i współpracę międzynarodową — oraz kierunki interwencji, w tym migrację do kryptografii postkwantowej. Punkt odniesienia dla wymagań bezpieczeństwa stawianych dostawcom systemów AI dla sektora publicznego

REGULACJE I NORMY

Superteam *Superzespół człowiek-AI*

Zespół, w którym agent AI wzmacnia możliwości pracownika, a człowiek zapewnia osąd, odpowiedzialność i znajomość kontekstu. Cel wdrożenia to nie zastąpienie ludzi, lecz zespoły, które z AI robią więcej i lepiej

WDROŻENIE I ZARZĄDZANIE

System Availability *Dostępność systemu AI (system availability)*

Odsetek czasu, w którym system AI jest dostępny i działa zgodnie z uzgodnionym poziomem usług; typowy cel dla systemów produkcyjnych to co najmniej 99,9% w skali miesiąca. Zob.

Uptime

SUWERENNOŚĆ I ARCHITEKTURA

System S46

System teleinformatyczny Krajowego Systemu Cyberbezpieczeństwa (art. 46 ustawy o KSC), prowadzony przez ministra właściwego do spraw informatyzacji i utrzymywany przez NASK, służący do zgłaszania i obsługi incydentów, wymiany informacji o zagrożeniach oraz prowadzenia wykazu podmiotów kluczowych i ważnych. Procedury reagowania na incydenty w systemach AI powinny być z nim zgodne

REGULACJE I NORMY · W regulacji: NIS2/uKSC

T

Taksonomia źródeł CDF *CDF Source Taxonomy – klasy W, M, N, D*

Czteroklasowa klasyfikacja źródeł kanonu CDF, wprowadzona w 1.6.0. Klasa W: akty prawa powszechnie obowiązujące lub w okresie dostosowawczym; wymagania obligatoryjne i bramkowane fazowo, stan stosowania w kolumnie Status matrycy E.1. Klasa M: dokumenty organów publicznych i regulatorów bez mocy wiążącej; domyślnie włączone dla objętych profili,

wylączenie wymaga uzasadnienia w Compliance Gap Analysis. Klasa N: normy i standardy techniczne; rama dojrzałości bez domniemania zgodności do czasu przywołania w Dzienniku Urzędowym UE. Klasa D: badania i dane empiryczne; zasilają Evidence Book i kalibrują wartości domyślne profili ACE.

METODYKA CDF

TCO *TCO — całkowity koszt posiadania*

Całkowity koszt posiadania rozwiązania w całym cyklu życia — nie tylko zakup lub budowa, lecz także integracja, utrzymanie, energia, personel, zgodność regulacyjna i koszt ewentualnej zmiany dostawcy. Dla systemów AI liczony zwykle w horyzoncie 3–5 lat

WDROŻENIE I ZARZĄDZANIE

TK Pending (flaga) *Constitutional Tribunal Pending (flaga CDF)*

Oznaczenie pozycji regulacyjnej objętej zawisłym postępowaniem przed Trybunałem Konstytucyjnym. Flaga oznacza ryzyko prawne wynikające z zawisłego postępowania, nie zaś zawieszenie stosowania przepisów; przy kontroli następczej ustawa obowiązuje do czasu ewentualnego orzeczenia o niezgodności. Stan na 21.09.2026: wniosek Prezydenta RP z 19.02.2026 o kontrolę następczą nowelizacji uKSC (Dz.U. 2026 poz. 252; zaskarżone m.in. art. 67b–67e o dostawcach wysokiego ryzyka), zarejestrowany w TK w sierpniu 2026 pod sygn. K 19/26 (wg źródeł wtórnych – potwierdzić na trybunal.gov.pl). Ustawa obowiązuje od 3.04.2026.

METODYKA CDF

Token *Jednostka tekstu*

Podstawowa jednostka, na jaką model językowy dzieli tekst — fragment słowa, całe słowo albo znak, zależnie od tokenizera. W tokenach liczy się długość promptu, odpowiedzi i koszt użycia modelu; w polszczyźnie jedno słowo to zwykle 2–3 tokeny

PODSTAWY AI

Token Budget Governance *Zarządzanie budżetem tokenów*

Mechanizm kontroli kosztów operacyjnych agentów AI przez limity zużycia tokenów modelu językowego — na agenta, na interakcję i na okres rozliczeniowy — z alertami i automatycznym zatrzymaniem po przekroczeniu. Szczególnie ważny w środowiskach, gdzie agenci wywołują się nawzajem i koszty mogą rosnąć lawinowo

WDROŻENIE I ZARZĄDZANIE

Transparency Notice (Powiadomienie o przejrzystości) *Obowiązek*

przejrzystości wobec użytkownika (art. 50 AI Act)

Obowiązek poinformowania osoby, że wchodzi w interakcję z systemem AI, oraz oznaczania treści wygenerowanych lub zmanipulowanych przez AI (art. 50 AI Act, stosowany od 02.08.2026). Dotyczy chatbotów, deepfake’ów, syntetycznych obrazów, dźwięku i tekstu oraz systemów rozpoznawania emocji i kategoryzacji biometrycznej

REGULACJE I NORMY · W regulacji: AI Act

Triada wydania CDF *CDF Release Triad*

Trzy zsynchronizowane artefakty publikowane w każdym wydaniu CDF: (1) Metodyka DOCX (PL), (2) Słownik DOCX, (3) Aplikacja webowa. Wszystkie trzy są ciągle doskonalone i muszą być spójne w każdym wydaniu. Źródło: proces wydawniczy CDF. Wprowadzone w v1.4.4.

METODYKA CDF

U

Układ z Komisją *Układ z Komisją (nadzwyczajne złagodzenie sankcji)*

Procedura ugody administracyjnej przewidziana w ustawie o systemach sztucznej inteligencji, w której podmiot naruszający przepisy może uzyskać od Komisji istotne obniżenie kary w zamian za współpracę, naprawienie szkody i wdrożenie działań naprawczych

REGULACJE I NORMY

Uptime *Dostępność (uptime)*

Procent czasu, w którym system pozostaje dostępny i działa zgodnie z oczekiwanym poziomem usług. Podstawowy wskaźnik w umowach SLA; 99,9% oznacza niecałe 45 minut niedostępności miesięcznie

SUWERENNOŚĆ I ARCHITEKTURA

US-EU AI Governance Divergence *Rozbieżności w regulacji AI między USA a UE*

Różnice między europejskim a amerykańskim podejściem do regulacji sztucznej inteligencji: UE ma jedno horyzontalne rozporządzenie oparte na klasyfikacji ryzyka, USA — przepisy sektorowe,

standardy dobrowolne (NIST) i regulacje stanowe. Dla firm działających na obu rynkach oznacza to różne obowiązki dokumentacyjne, inne definicje i odmienny reżim odpowiedzialności

REGULACJE I NORMY

Ustawa o systemach sztucznej inteligencji (SSI) *Ustawa z dnia 3 lipca 2026 r. o systemach sztucznej inteligencji*

Krajowa ustawa zapewniająca stosowanie AI Act w Polsce (Dz.U. z 2026 r. poz. 1003).

Ustanawia Komisję Rozwoju i Bezpieczeństwa Sztucznej Inteligencji jako organ nadzoru rynku, reguluje opinie indywidualne, kontrole i postępowania, układ w sprawie złagodzenia sankcji, jednostki notyfikowane, piaskownicę regulacyjną oraz kary. Weszła w życie 11.08.2026; przepisy o opiniach, kontroli, postępowaniu, układzie i karach stosuje się od 28.10.2026

REGULACJE I NORMY · W regulacji: AI Act

V

Vendor Disruption Protocol (VDP) *Protokół na wypadek utraty dostawcy (VDP)*

Uporządkowany plan postępowania na wypadek utraty dostępu do dostawcy AI — wskutek sankcji, wpisania na listę podmiotów wykluczonych, upadłości lub zmiany warunków. Obejmuje zapasowego dostawcę, możliwość przejścia na model suwerenny, procedurę migracji, ocenę skutków w łańcuchu dostaw i zapisy o odpowiedzialności wykonawcy

SUWERENNOŚĆ I ARCHITEKTURA

Vendor Lock-in *Uzależnienie od dostawcy (vendor lock-in)*

Sytuacja, w której organizacja ma istotnie ograniczoną możliwość zmiany dostawcy technologii — z powodu formatu danych, architektury, warunków umowy lub braku kontroli nad kodem i integracjami. W AI dotyczy też zależności od konkretnego modelu i jego promptów. Prawo do zmiany dostawcy usług chmurowych gwarantuje Data Act

SUWERENNOŚĆ I ARCHITEKTURA · W regulacji: Data Act

Vendor Risk Assessment *Ocena ryzyka dostawcy AI*

Proces oceny ryzyka związanego z dostawcami modeli AI i komponentów infrastruktury. Obejmuje kryteria: jurysdykcja danych, przejrzystość łańcucha dostaw, gwarancje ciągłości, możliwość audytu, zgodność z AI Act oraz strategia wyjścia. Wymagany m.in. przez DORA i NIS2 wobec dostawców ICT

SUWERENNOŚĆ I ARCHITEKTURA · W regulacji: AI Act · DORA · NIS2/uKSC

Vulnerability Disclosure (Zgłaszanie podatności) *Zgłaszanie i skoordynowane ujawnianie podatności*

Obowiązek producenta produktu z elementami cyfrowymi do obsługi podatności przez cały okres wsparcia — w tym polityki skoordynowanego ujawniania (Cyber Resilience Act, art. 13 i załącznik I) — oraz zgłaszania aktywnie wykorzystywanych podatności i poważnych incydentów do właściwego CSIRT i ENISA (art. 14). Dotyczy także komponentów AI wbudowanych w produkt

REGULACJE I NORMY · W regulacji: CRA

W

Watermarking (AI) *Znakowanie treści generowanych przez AI*

Oznaczanie treści generowanych przez systemy AI tak, by dało się je rozpoznać: w warstwie maszynowej (metadane, niewidoczny znak wodny) i w razie potrzeby w warstwie widocznej. Obowiązek wynika z art. 50 AI Act, stosowanego od 2 sierpnia 2026; dla systemów wprowadzonych do obrotu przed tą datą przewidziano okres przejściowy

PODSTAWY AI · W regulacji: AI Act

Wejście w życie ustawy o SSI *Harmonogram stosowania ustawy o systemach sztucznej inteligencji*

Trzon ustawy o systemach sztucznej inteligencji obowiązuje od 11.08.2026 (14 dni od ogłoszenia). Przepisy o opiniach indywidualnych oraz rozdziały o kontroli, postępowaniu, układzie, karach administracyjnych i przepisach karnych stosuje się od 28.10.2026 — od tej daty działa pełny reżim kontrolno-sankcyjny

REGULACJE I NORMY · W regulacji: ustawa o systemach SI

Wiedza Ukryta *Tacit Knowledge — wiedza niejawna, nieudokumentowana*

Wiedza, którą pracownik ma, ale która nie jest zapisana w systemach organizacji — w głowie, na dysku lokalnym, w prywatnych notatkach i nieformalnych interpretacjach. Jej pozyskanie i

uporządkowanie jest celem Cyfrowego Bliźniaka Poznawczego; bez tego asystent AI zna procedury, ale nie zna praktyki

PODSTAWY AI

Workflow Redesign Mandate *Zasada przeprojektowania procesu przed automatyzacją*

Zasada, zgodnie z którą każdy proces musi zostać przeprojektowany, zanim zintegruje się z nim agenta AI. Sztucznej inteligencji nie nakłada się na nieefektywne procesy z przeszłości — inaczej automatyzuje się i powiela ich wady

WDROŻENIE I ZARZĄDZANIE

Wyrób medyczny AI (URPL) *Wyrób medyczny z systemem AI (URPL)*

Wyrób medyczny, którego elementem bezpieczeństwa jest system AI lub który sam jest systemem AI (MDR 2017/745, IVDR 2017/746). Podlega ocenie zgodności z udziałem jednostki notyfikowanej według obu reżimów; w Polsce Urząd Rejestracji Produktów Leczniczych, Wyrobów Medycznych i Produktów Biobójczych współpracuje z organem notyfikującym w rozumieniu ustawy o systemach SI

REGULACJE I NORMY · W regulacji: ustawa o systemach SI · MDR/IVDR

X

XAI *Explainable AI — wyjaśnialna sztuczna inteligencja*

Techniki i metody pozwalające zrozumieć, jak model doszedł do decyzji, na jakiej podstawie i z jakim poziomem pewności. W systemach wysokiego ryzyka AI Act wymaga, by decyzje dało się wyjaśnić osobom nadzorującym (art. 13) i osobom, których dotyczą (art. 86)

PODSTAWY AI · W regulacji: AI Act

Z

Zero Trust *Architektura zerowego zaufania (zero trust)*

Model bezpieczeństwa zakładający, że żaden użytkownik, urządzenie ani agent nie jest domyślnie zaufany, niezależnie od położenia w sieci — każdy dostęp jest uwierzytelniany, autoryzowany i rejestrowany. W środowiskach wieloagentowych oznacza osobną tożsamość i minimalne uprawnienia dla każdego agenta

SUWERENNOŚĆ I ARCHITEKTURA

Znacząca modyfikacja *Istotna zmiana systemu AI (substantial modification)*

Zmiana systemu AI po wprowadzeniu do obrotu, nieprzewidziana w pierwotnej ocenie zgodności — np. dostrojenie modelu, ciągłe uczenie lub zmiana przeznaczenia. W rozumieniu AI Act (art. 3 pkt 23, art. 25) może uczynić podmiot stosujący dostawcą, a w rozumieniu dyrektywy 2024/2853 (art. 8 ust. 2) — producentem odpowiedzialnym za produkt. Wymaga rejestracji i ponownej oceny

REGULACJE I NORMY · W regulacji: AI Act · PLD 2024/2853