

AI jako wyzwanie strategiczne

Przewodnik dla rad nadzorczych

Opracowane przez Stowarzyszenie Niezależnych Członków
Rad Nadzorczych (SNCRN)

Autorzy:

Katarzyna Kazior
Michał Lasocki



Spis treści

1.	STRESZCZENIE MENEDŻERSKIE	4
1.1.	Kontekst globalny i europejski.....	4
1.2.	Jak rady nadzorcze powinny podejść do AI	4
1.3.	Trendy i odpowiedzialność nadzorcza	4
1.4.	Praktyczne działania dla rad nadzorczych	5
2.	KONTEKST GEOPOLITYCZNY: EUROPA W OBLICZU AI – WYZWANIA I ODPOWIEDZI	6
2.1.	Streszczenie menedżerskie kontekstu globalnego	6
2.2.	Firmy polskie i europejskie mogą skutecznie konkurować w dziedzinie AI. Oto warunki	7
2.3.	Raport Dragiego i jego znaczenie dla Europy	7
2.4.	Prymat USA w dziedzinie AI	10
2.5.	Sukcesy Chin w dziedzinie AI	12
2.6.	Globalne wnioski o charakterze uniwersalnym	14
2.7.	Możliwa odpowiedź firm europejskich	15
2.8.	Bezpieczeństwo geopolityczne	17
2.9.	Sytuacja Polski i Europy Środkowo-Wschodniej	18
2.10.	Podsumowanie sytuacji geopolitycznej	18
3.	JAK PRZYGOTOWAĆ SIĘ DO NADZORU NAD AI?	19
3.1.	Podstawowe kwestie AI w spółkach – co trzeba wiedzieć?	19
3.2.	Jakie trendy przeważają w AI w 2025 roku?	20
3.3.	GenAI vs Data AI – znaczenie danych	22
3.4.	Skąd się uczyć o AI? Szkolenie jako całościowa odpowiedzialność rad nadzorczych i kadry	24
4.	FUNDAMENTY EFEKTYWNEGO NADZORU NAD AI W ORGANIZACJI	26
4.1.	Istotność strategiczna i ryzyko braku działania	26
4.1.1.	Bilans otwarcia AI w organizacji	26
4.1.2.	Umiejscowienie odpowiedzialności	27
4.1.3.	Osadzenie AI w strategii organizacji	27
4.1.4.	Identyfikacja możliwości biznesowych	28
4.1.5.	Ekonomiczna analiza kosztów i korzyści	28
4.1.6.	Zarządzanie talentem i edukacja	29
4.1.7.	Realistyczne osadzenie w możliwościach spółki	29
4.1.8.	Jakie zadania zlecamy podmiotom zewnętrznym, a co pozostaje wewnątrz organizacji (czyli build vs. buy)	30
4.2.	Odpowiedzialne AI	31
4.2.1.	Zgodność z wartościami firmy	31
4.2.2.	Integracja z systemami kontroli wewnętrznej	31
4.2.3.	Polityka oceny zaufania i wiarygodności	32
4.2.4.	Cyberbezpieczeństwo w kontekście AI	32
4.2.5.	Nieautoryzowane użycie AI przez pracowników	33
4.2.6.	Przeciwdziałanie stronniczości algorytmów	33
4.2.7.	Ochrona danych i prywatności	34
4.2.8.	Polityki wewnętrzne	34
4.2.9.	Wyjaśnialność i transparentność	34
4.2.10.	Podsumowanie	35
4.3.	Ryzyka i Regulacje – monitorowanie otoczenia regulacyjnego	36
4.3.1.	Przegląd obecnych regulacji dotyczących AI	36
4.3.2.	Co jest chronione prawem, a co nie?	40
4.3.3.	Sytuacja w Polsce – wdrażanie regulacji i programy wsparcia	47
5.	CO DALEJ? ZINTEGROWANIE TEMATU AI W BIEŻĄCYM PROCESIE NADZORU	48
5.1.	Normalizacja kwestii AI	48
5.2.	Dostosowanie struktury rady nadzorczej	48

5.3.	Ocena kadry menedżerskiej w spółce	49
5.4.	Planowanie budżetu i planów na długi termin	49
5.5.	Ofensywa zamiast defensywy – co dla każdego	50
5.6.	Ciągłe uczenie się – jak to może wyglądać	50
6.	POZIOMY INTEGRACJI AI ZE STRATEGIĄ FIRMY – PRZYKŁADY MOŻLIWYCH OBSZARÓW W FIRMACH ORAZ STOPNI INTEGRACJI	51
6.1.	Poziom 1 – wdrażanie AI w codziennych zadaniach	51
6.2.	Poziom 2 – kolejny krok to modyfikacja kluczowych procesów.	53
6.3.	Poziom 3. Najbardziej zaawansowane wykorzystanie AI to tworzenie nowych produktów i usług	54
6.4.	Przykłady wpływu na wybrane branże AI połączonego z nowoczesnymi urządzeniami. Szanse i zagrożenia	55
7.	O AUTORACH:	65
8.	LINKI DO STRON WYKORZYSTANYCH W NINIEJSZYM DOKUMENCIE.....	66

1. Streszczenie Menedżerskie

Rozwój sztucznej inteligencji (AI) jest jednym z kluczowych czynników kształtujących geopolitykę, konkurencyjność gospodarczą oraz strategię firm. Rady nadzorcze w Europie, w tym w Polsce, muszą traktować AI jako priorytet strategiczny – zarówno z perspektywy szans, jak i ryzyk.

1.1. Kontekst globalny i europejski

- **AI to pole globalnej rywalizacji** między USA, Chinami i – wciąż niedostatecznie obecnie – Europą. Europa stoi przed wyzwaniem budowy suwerenności technologicznej.
- **Raport Dragiego** wskazuje na potrzebę masywnych i skoordynowanych inwestycji w AI, infrastrukturę cyfrową i talenty. Europa musi działać szybko, by nie pozostać na marginesie.
- **Polska i Europa Środkowo-Wschodnia** dysponują talentami i kompetencjami cyfrowymi, ale brakuje im skali finansowej i danych potrzebnych do globalnych wdrożeń AI.

1.2. Jak rady nadzorcze powinny podejść do AI

- **Zrozumienie strategicznego znaczenia AI** – AI może zwiększać wydajność, umożliwiać nowe modele biznesowe, ale jej niewłaściwe użycie może prowadzić do strat reputacyjnych i prawnych.
- **Integracja AI z celami firmy** – AI nie może być projektem technologicznym oderwanym od strategii. Powinna być częścią długofalowej wizji rozwoju.
- **Identyfikacja obszarów wartości biznesowej** – kluczowe zastosowania AI to automatyzacja zadań, optymalizacja procesów i tworzenie innowacyjnych produktów.
- **Ocena ryzyk AI** – prawne, technologiczne, organizacyjne i społeczne (np. uprzedzenia w danych). Należy zadbać o zgodność z regulacjami (np. AI Act, DORA, NIS2).

1.3. Trendy i odpowiedzialność nadzorcza

- **2025 to rok regulacji AI** – AI Act i inne przepisy wymagają od firm odpowiednich działań compliance.
- **Explainability i etyka AI** – coraz większy nacisk kładzie się na wyjaśnialność algorytmów, etykę, bezpieczeństwo i ochronę danych.
- **Wzrost znaczenia GenAI i AI Agents** – rady muszą rozróżniać pomiędzy GenAI (np. ChatGPT) a Data AI (np. systemy predykcyjne).

1.4. Praktyczne działania dla rad nadzorczych

- **Audyt stanu AI w firmie** i mapa drogowa dalszego rozwoju.
- **Wyznaczenie odpowiedzialności** (np. CDO, dedykowane komitety).
- **Szkolenia dla rady** – AI to temat wymagający ciągłej nauki i aktualizacji wiedzy.
- **Monitoring otoczenia prawnego** – dynamiczne zmiany regulacyjne wymagają czujności.
- **Współpraca z ekspertami i uczelniami** – zewnętrzne wsparcie może znacząco zwiększyć dojrzałość AI w organizacji.

2. Kontekst geopolityczny: Europa w obliczu AI – wyzwania i odpowiedzi

2.1. Streszczenie menedżerskie kontekstu globalnego

Technologia AI stała się centralnym elementem globalnej rywalizacji geopolitycznej. Europa musi dążyć do strategicznej autonomii w tej dziedzinie, aby zachować suwerenność technologiczną i móc niezależnie kształtować swoje interesy gospodarcze i bezpieczeństwa.

Zrozumienie aktualnej pozycji Polski i Europy w kontekście globalnej rywalizacji AI jest kluczowe dla członków rad nadzorczych. AI nie jest już tylko technologiczną innowacją, staje się jednym z fundamentalnych czynników przewagi konkurencyjnej, bezpieczeństwa i rozwoju gospodarczego. Choć dominacja USA i Chin jest obecnie bezsporna, jak podkreślono w [raporcie Mario Draghiego](#): **Europa nadal ma szansę odegrać istotną rolę w tej transformacji.**

„Wyścig o przywództwo w AI nie jest jeszcze rozstrzygnięty. Europa ma wiele atutów, by odegrać istotną rolę w kształtowaniu przyszłości tej technologii, ale wymaga to odważnych decyzji, większych inwestycji i bardziej skoordynowanego podejścia, zarówno na poziomie instytucjonalnym, jak i korporacyjnym.”

„With the world now on the cusp of another digital revolution, triggered by the spread of artificial intelligence (AI), a window has opened for Europe to redress its failings in innovation and productivity and to restore its manufacturing potential.”

Raport Draghiego, część A str. 14

Raport ten, autorstwa byłego prezesa Europejskiego Banku Centralnego i premiera Włoch jest najbardziej kompleksową próbą osadzenia AI w szerszym kontekście konkurencyjności Europy. Przedstawia on nie tylko diagnozę opóźnień technologicznych UE, ale i konkretne zalecenia dla odzyskania pozycji w globalnym wyścigu. W części B znajduje się szczegółowe rozbieżności na sektory związane z branżą cyfrową, a w nim istotne stwierdzenie:

„UE musi mieć ambicję, aby stać się liderem w rozwoju sztucznej inteligencji w sektorach, w których ma przewagę, odzyskać i utrzymać kontrolę nad danymi oraz wrażliwymi usługami chmurowymi, a także zbudować solidny mechanizm finansowo-talentowy wspierający innowacje w zakresie obliczeń i AI.”

“ The EU must have the ambition to be a leader in developing AI for its sectors of strength, regain and retain control over data and sensitive cloud services, and develop a robust financial and talent flywheel to support innovation in computing and AI. ”

Raport Draghiego, część B str. 82

2.2. Firmy polskie i europejskie mogą skutecznie konkurować w dziedzinie AI. Oto warunki

Jeśli firmy europejskie połączą inwestycje, współpracę i unikalne wartości UE, możliwe jest konkurowanie w określonych niszach i obszarach. Odpowiedzialne i celowe wdrażanie AI może stać się fundamentem nowej, konkurencyjnej Europy.

Mimo wielu wyzwań, firmy zarówno z Polski jak i z całej Europy wciąż mają realną szansę konkurować w dziedzinie AI. Wymaga to jednak:

- zintensyfikowania inwestycji w talenty i infrastrukturę cyfrową,
- pogłębionej współpracy publiczno-prywatnej,
- koordynacji i współdziałania zamiast rozpraszania zasobów
- wykorzystania unikalnych atutów Europy, takich jak etyczne podejście, unikalne doświadczenia przemysłowe, transparentność i dbałość o prawa obywateli.

Odpowiedzialne i celowe wdrażanie AI może stać się fundamentem nowej, konkurencyjnej Europy. Sukces będzie zależał od odwagi decyzyjnej liderów – zarówno politycznych, jak i po stronie przedsiębiorstw – by na nowo zdefiniować rolę Europy w erze sztucznej inteligencji.

2.3. Raport Dragiego i jego znaczenie dla Europy

Raport Dragiego zawiera szereg jednoznacznych wskazówek dotyczących konieczności szybkiego i skoordynowanego działania Europy w obszarze nowych technologii – ze szczególnym naciskiem na AI.

Dokument ostrzega, że Europa może nie być w stanie realizować swoich ambicji gospodarczych, społecznych i strategicznych, jeśli nie zwiększy produktywności i nie zainwestuje w przełomowe technologie:

„Jeśli Europa nie stanie się bardziej produktywna, będziemy zmuszeni dokonać wyboru. Nie będziemy w stanie jednocześnie być liderem w nowych technologiach, wzorem odpowiedzialności klimatycznej i niezależnym graczem na arenie międzynarodowej. [...] To wyzwanie egzystencjalne.”

„If Europe cannot become more productive, we will be forced to choose. We will not be able to become, at once, a leader in new technologies, a beacon of climate responsibility and an independent player on the world stage. [...] This is an existential challenge.”

Raport Dragiego, Przedmowa

„Okolo 70% podstawowych modeli AI zostało opracowanych w Stanach Zjednoczonych od 2017 roku, a zaledwie trzech amerykańskich ‘hiperskalowych’ dostawców usług chmurowych kontroluje ponad 65% globalnego – i europejskiego – rynku chmury. Największy europejski operator chmurowy posiada jedynie 2% udziału w rynku UE.”

„Around 70% of foundational AI models have been developed in the US since 2017, and just three US ‘hyperscalers’ account for over 65% of the global – and European – cloud market. The largest European cloud operator accounts for just 2% of the EU market.”

Raport Draghiego, część A, rozdział 2

Brak zdecydowanych działań pogłębia lukę inwestycyjną – europejskie przedsiębiorstwa wydają na badania i rozwój o 270 miliardów euro mniej niż ich amerykańskie odpowiedniki. W konkluzji raport podkreśla konieczność stworzenia nowej architektury polityki przemysłowej i innowacyjnej, która pozwoli Europie nie tylko nadążyć za światową czołówką, ale także tworzyć własne przewagi technologiczne w oparciu o wartości, talenty i strategiczne zasoby.

Główne zalecenia raportu obejmują:

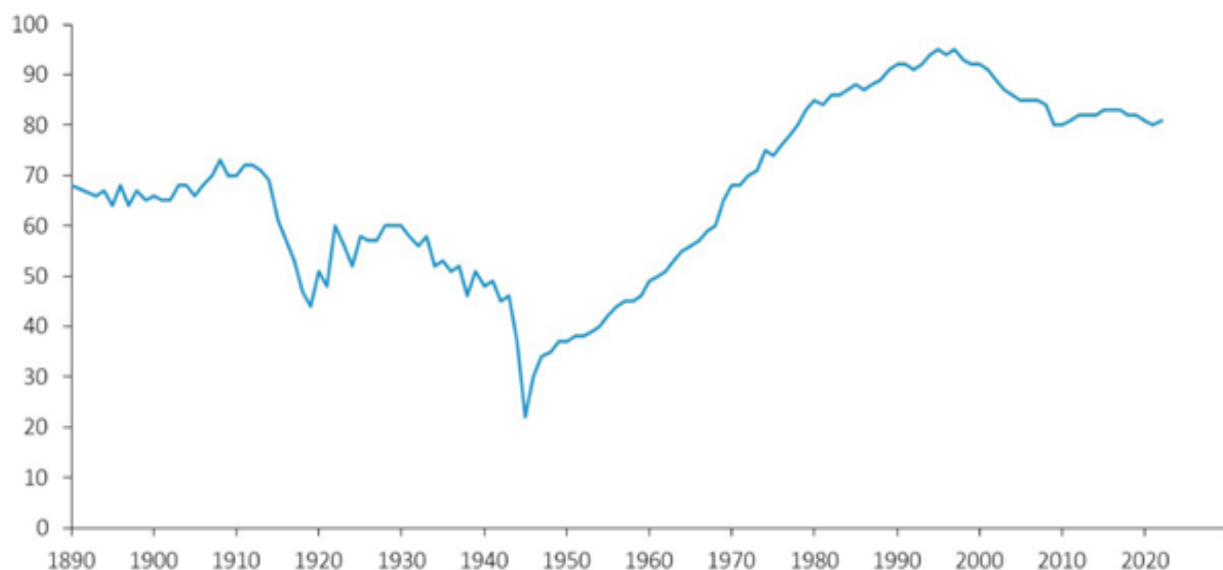
- Potrzebę stworzenia wspólnej strategii inwestycyjnej w badania i rozwój (B+R), zwłaszcza w AI, półprzewodniki i technologie obliczeniowe,
- Konieczność zwiększenia inwestycji publicznych i prywatnych w AI,
- Istotność elastycznego i przewidywalnego otoczenia regulacyjnego,
- Konieczność integracji rynku cyfrowego UE jako warunku skalowalności technologii,
- Konsolidację europejskiego podejścia do badań oraz rewizję ram regulacyjnych.

Omawiany raport ilustruje spadek produktywności w Europie. Nasz kontynent potrzebuje szybszego wzrostu produktywności, aby utrzymać zrównoważone tempo rozwoju w obliczu niekorzystnych trendów demograficznych. Wydajność pracy w UE zbliżyła się z poziomu 22% amerykańskiego poziomu w 1945 roku do 95% w 1995 roku, ale później tempo wzrostu produktywności pracy w Europie spadło bardziej niż w USA i ponownie obniżyło się poniżej 80% poziomu amerykańskiego. Poniższy wykres ilustruje zmiany względnej produktywności UE w odniesieniu do USA.

WYKRES 1

Porównanie produktywności UE do USA 1890–2022

Indeks (USA=100)



Uwaga: Dane dla UE są reprezentowane pośrednio poprzez cofnięcie w czasie danych z rachunków narodowych Niemiec, Francji, Włoch, Hiszpanii, Niderlandów, Belgii, Irlandii, Austrii, Portugalii, Finlandii i Grecji. Do opracowania danych dotyczących produktywności pracy wykorzystano pięć różnych serii: PKB, zasób kapitału, zatrudnienie, średnią liczbę przepracowanych godzin oraz populację. Zasób kapitału jest tworzony na podstawie dwóch serii danych o inwestycjach – w budownictwo i w sprzęt. Inwestycje oraz PKB są przedstawione w ujęciu wolumenowym i w narodowej walucie z 2010 roku, a następnie przeliczane na dolary z 2010 roku przy użyciu współczynnika konwersji parytetu siły nabywczej (PPP).

Note: The EU is proxied by backdating national accounting data from Germany, France, Italy, Spain, the Netherlands, Belgium, Ireland, Austria, Portugal, Finland and Greece. To build the labour productivity data, five different series were used: GDP, capital stock, employment, average hours worked, and population. Capital stock is built using two series of investment – construction and equipment. Investment and GDP are taken in volume and in national currency of 2010, they are then turned into \$2010 using a ppp conversion rate

Konkluzje raportu wskazują, że obecne środowisko regulacyjne UE nie sprzyja innowacjom w zakresie AI w porównaniu z innymi rynkami na świecie. Potrzebujemy bardziej zwinnego podejścia, które będzie chronić europejskie wartości, jednocześnie umożliwiając firmom szybsze wdrażanie nowych rozwiązań. Czynniki te, w połączeniu z trudnościami w pozyskaniu finansowania na duże i ambitne projekty na późniejszych etapach ich rozwoju, sprawiają, że w Europie brakuje championów technologicznych nawet w obszarach, gdzie firmy mają doświadczenie:

“Europa zajmuje silną pozycję w dziedzinie robotyki autonomicznej, skupiając około 22% globalnej aktywności oraz w usługach opartych na sztucznej inteligencji – z udziałem około 17%. Jednak innowacyjne firmy cyfrowe w Europie zazwyczaj nie są w stanie się skalować i przyciągać finansowania, co odzwierciedla ogromna luka w finansowaniu na późniejszych etapach rozwoju między UE a USA. W rzeczywistości, w ciągu ostatnich pięćdziesięciu lat w UE nie powstała żadna firma, której kapitalizacja rynkowa przekroczyłaby 100 miliardów euro, podczas gdy w USA wszystkie sześć firm wycenianych obecnie na ponad 1 bilion dolarów powstało właśnie w tym okresie.”

“Europe holds a strong position in autonomous robotics, hosting around 22% of worldwide activity, and in AI services, hosting around 17% of activity. But innovative digital companies are generally failing to scale up in Europe and attract finance, reflected in a huge gap in later-stage financing between the EU and the US. In fact, there is no EU company with a market capitalisation over

EUR 100 billion that has been set up from scratch in the last fifty years, while in the US all six companies with a valuation above EUR 1 trillion have been created over this period”

Raport Draghiego, str. 24

Według Europejskiego Banku Inwestycyjnego (2023), luka inwestycyjna w dziedzinie AI między Europą a USA poszerza się, co potwierdza pilną potrzebę działania.

„Dane z badania inwestycyjnego EBI (EIB Investment Survey – EIBIS) pokazują również, że firmy z UE rzadziej niż firmy z USA wprowadzają innowacje lub wdrażają nowe technologie. W 2022 roku ta różnica powiększyła się o około 10 punktów procentowych, osiągając łącznie 19 punktów procentowych. Luka ta wynika głównie z rzadszych inwestycji firm unijnych we wdrożenie nowych technologii i praktyk. „

„Data from the EIB Investment Survey (EIBIS) also show that EU firms are less likely to innovate or to adopt new technologies than US firms. The gap actually widened by around 10 percentage points in the 2022 survey, to 19 percentage points. This gap is driven by less frequent investment by EU firms in the adoption of new technologies and practices.”

wdrożenie nowych technologii i praktyk.”

Economic Investment Report, str. 2

Dane z raportu CB Insights State of AI dotyczące pierwszego kwartału 2025 ukazują łączne sumy inwestycji w AI w obszarach private equity, venture capital, oraz seed :

- USA – 60,9 mld USD
- Europa – 3,5 mld USD (W Europie siedemnastokrotnie mniej niż w USA)
- Azja (w tym Chiny) nominalnie 1,5 mld USD. Jest to kwota wielokrotnie zaniżona, ze względu na to, że AI w Chinach rozwija się w dużej mierze „w obrębie istniejących dużych przedsiębiorstw cyfrowych” a wspomniany raport dotyczy aktywności inwestorów kapitałowych. Osiągnięcia Chin nie były by możliwe przy tak niskim wykazanym poziomie inwestycji. W dalszej części dokumentu opisano narodowy program inwestycyjny Chin warty dziesiątki mld USD

2.4. Prymat USA w dziedzinie AI

W obszarze AI dominują dwie potęgi – **Stany Zjednoczone i Chiny**. [Raport Stanford HAI](#) oraz OECD AI Policy Observatory wskazują na konkretne przyczyny przewagi tych krajów:

Stany Zjednoczone przewodzą pod względem liczby firm AI (ponad 13 tys. startupów), inwestycji venture capital (63% globalnego udziału w 2023 roku) oraz zatrudniania talentów AI z całego świata. Sektor prywatny USA zainwestował w 2023 roku ponad 120 miliardów dolarów w badania i rozwój AI. ([Stanford AI Index](#), 2024).

8 z 10 największych firm technologicznych rozwijających AI pochodzi z USA:

OpenAI to firma znana jako twórca popularnego rozwiązania **ChatGPT** oraz wielu innych. Firma rozwija szereg zaawansowanych modeli, takich jak **GPT-4**, **Codex** czy **DALL-E** (wielu członków zespołu OpenAI to Polacy, choć na razie na potrzeby naszej gospodarki niewiele z tego wyniknęło).

Google to firma znana jako twórca całego ekosystemu IT wywodzącego się z narzędzi Google. W skład tego ekosystemu wchodzi m.in.: wyszukiwarka Google (ma zindeksowane zasoby internetu), przeglądarka internetowa **Chrome**, usługa pocztowa **Gmail**, nawigacja **Google Maps**, system reklam internetowych **Google AdWords**. Oprócz tu wymienionych Google oferuje setki usług w dziedzinie informatycznej. Dostęp do olbrzymich zbiorów danych znacznie ułatwił firmie stworzenie bardzo zaawansowanych modeli sztucznej inteligencji znanych obecnie pod nazwą **Gemini**, kontynuując rozwój wcześniejszych technologii takich jak np. **TensorFlow** czy **BERT**.

Microsoft to jeden z głównych globalnych graczy w dziedzinie nowych technologii i sztucznej inteligencji. Tworzy własny ekosystem usług, obejmujący wyszukiwarkę i przeglądarkę **Bing**, pakiet biurowy w wersji chmurowej **Microsoft 365**, a także platformę chmurową **Azure**, która stała się jednym z filarów przetwarzania danych i trenowania dużych modeli językowych. Microsoft jest strategicznym partnerem **OpenAI** – to właśnie dzięki tej współpracy integruje modele GPT (w tym GPT-4) z produktami takimi jak **Copilot** w **Word**, **Excel**, **Outlook** czy **Bing Chat**. Azure oferuje także usługi AI jako gotowe API lub możliwość trenowania własnych modeli, co czyni go kluczowym środowiskiem dla wdrożeń korporacyjnych.

Meta (wcześniej Facebook) to właściciel globalnych platform społecznościowych, takich jak **Facebook**, **Instagram** czy **WhatsApp**, ale również jeden z najbardziej aktywnych ośrodków badań nad AI na świecie. Firma stworzyła i udostępniła rodzinę otwartych modeli językowych **LLaMA (Large Language Model Meta AI)** – od wersji 1 aż po najnowszą **LLaMA 3**, która konkuruje z GPT-4 pod względem jakości. Meta prowadzi także prace nad multimodalnymi systemami AI oraz generatywnymi modelami obrazów i wideo. Otwarty charakter większości ich modeli sprawia, że są one szeroko wykorzystywane w badaniach, edukacji i przemyśle, a także stanowią realną alternatywę dla zamkniętych rozwiązań komercyjnych.

NVIDIA: Choć NVIDIA zaczęła jako producent kart graficznych (GPU), dziś jest kluczowym graczem w rozwoju sztucznej inteligencji. Ich procesory graficzne napędzają trenowanie największych modeli AI na świecie – w organizacjach takich jak OpenAI, Google, czy Meta. NVIDIA dostarcza nie tylko sprzęt, ale także pełen ekosystem oprogramowania AI (np. **CUDA**, **TensorRT**, **NeMo**).

Anthropic to amerykańska firma specjalizująca się w sztucznej inteligencji, założona przez byłych pracowników OpenAI. Zyskała rozgłos dzięki rozwojowi własnych modeli językowych z serii Claude – alternatywy dla GPT, deklarując jej zaprojektowanie z myślą o większym bezpieczeństwie, przewidywalności i zgodności z wartościami użytkowników.

Amazon to nie tylko wielki gracz globalnego e-commerce, ale także jeden z filarów infrastruktury cyfrowej dzięki swojej platformie **AWS (Amazon Web Services)** – największej chmurze obliczeniowej na świecie. Amazon udostępnia własne rozwiązania AI, w tym model językowy **Titan**, narzędzie do tworzenia chatbotów **Bedrock**, a także technologie rozpoznawania obrazu i mowy. AWS dostarcza infrastrukturę dla wielu firm budujących własne modele AI, co czyni Amazon istotnym ogniwem w łańcuchu wartości sztucznej inteligencji.

X (dawniej **Twitter**), dysponując jedną z największych na świecie baz danych językowych w postaci miliardów publicznych tweetów, rozwija własny model językowy o nazwie **Grok**. Model ten ma być ściśle zintegrowany z platformą społecznościową i oferować odpowiedzi w czasie rzeczywistym w ramach

abonamentu **X Premium**. Projekt Grok jest rozwijany przez spółkę zależną **xAI**, która deklaruje ambicję stworzenia AI zdolnej do „racjonalnego rozumowania”.

65%

najlepszych światowych talentów w dziedzinie AI przyciągają amerykańskie uczelnie ([McKinsey Global Institute](#), 2023).

Firmy amerykańskie posiadają dostęp do ogromnych zbiorów danych poprzez dominujące platformy cyfrowe.

Jak zauważa [raport Goldman Sachs](#), przewaga USA w dziedzinie AI wynika z unikalnego ekosystemu łączącego kapitał venture, wiodące uniwersytety i talenty, wspieranego przez przyjazne regulacje i masywne zbiory danych generowane przez platformy cyfrowe. Bank ten przygotował kilka interesujących raportów na temat wpływu AI na gospodarkę.

Wg ekspertów z tego banku, sama generatywna sztuczna inteligencja może zwiększyć globalny PKB o 7 bilionów dolarów w ciągu dekady i podnieść roczny wzrost produktywności o 1,5 punktu procentowego. AI ma potencjał do automatyzacji 25% zadań w krajach rozwiniętych, ale jej wdrożenie napotyka bariery, takie jak wysokie koszty, ograniczony obszar wdrożenia, poziom penetracji rynkowej, zakres stosowania i spopularyzowania i brak odpowiednio przeszkolonej kadry.

2.5. Sukcesy Chin w dziedzinie AI

Chiny koncentrują się na integracji AI z państwowymi programami gospodarczymi. Dzięki temu mogą skoordynować działania i uniknąć rozpraszania zasobów. Zagadnienia związane z AI są opisane w [Narodowej strategii rozwoju AI](#) z 2017 roku z planem inwestycyjnym wartym ponad 150 miliardów dolarów USA. Firmy chińskie posiadają również dostęp do ogromnych zbiorów danych dzięki dużej populacji i mniej restrykcyjnemu podejściu do prywatności. Państwo chińskie dostarcza systematyczne wsparcie państwa dla „narodowych mistrzów” w dziedzinie AI.

Baidu to chiński gigant technologiczny, który rozwija własny pełen ekosystem sztucznej inteligencji – od infrastruktury sprzętowej po oprogramowanie i modele językowe. Firma opracowała serię dużych modeli językowych o nazwie **ERNIE (Enhanced Representation through Knowledge Integration)**, konkurującą z rozwiązaniami typu GPT. Baidu posiada również własne układy AI – **Kunlun**, zaprojektowane do obsługi zadań uczenia maszynowego i przetwarzania języka naturalnego. W zakresie oprogramowania Baidu rozwija platformę **PaddlePaddle** – otwarte środowisko uczenia głębokiego, będące chińską alternatywą dla TensorFlow czy PyTorch. Firma integruje AI m.in. z wyszukiwarką, usługami chmurowymi, systemem rozpoznawania mowy i autonomicznymi pojazdami.

Alibaba to nie tylko wielki gracz w e-commerce i właściciel Aliexpress, ale również jeden z głównych dostawców rozwiązań AI w Chinach. Firma oferuje własną chmurę obliczeniową **Alibaba Cloud**, w

ramach której rozwija pakiet narzędzi AI, w tym **ModelScope** – otwartą platformę z setkami gotowych modeli językowych, wizualnych i multimodalnych. Flagowy model językowy Alibaby to **Tongyi Qianwen**, który został zintegrowany z ekosystemem usług takich jak przeglądarka, chatboty, usługi finansowe i logistyka. Alibaba projektuje także własne układy AI – **Hanguang**, zoptymalizowane do przetwarzania obrazów i tekstu w chmurze.

Tencent, właściciel platform takich jak **WeChat** i **QQ**, inwestuje intensywnie w rozwój własnych modeli AI, skupiając się na ich wdrażaniu w usługach konsumenckich i biznesowych. Firma rozwija model językowy **Hunyuan**, przeznaczony do przetwarzania języka naturalnego, generowania treści i wspomagania pracy programistów. W obszarze infrastruktury **Tencent Cloud** oferuje zestaw usług AI, w tym narzędzia do rozpoznawania obrazu, mowy i przetwarzania dokumentów. Firma rozwija również własny framework AI oraz środowiska uczenia głębokiego, dopasowane do chińskiego rynku i lokalnych zastosowań.

Huawei to jeden z filarów chińskiej autonomii technologicznej w dziedzinie AI. Firma rozwija zarówno **sprzęt**, jak i **pełen ekosystem oprogramowania**. Kluczowym elementem są procesory **Ascend** – układy AI własnej produkcji, zaprojektowane do obsługi dużych modeli językowych i wizualnych. Na poziomie oprogramowania Huawei oferuje otwartą platformę **MindSpore**, alternatywę dla TensorFlow, zoptymalizowaną pod architekturę Ascend. AI od Huawei znajduje zastosowanie w systemach nadzoru, telekomunikacji 5G, inteligentnych miastach, centrach danych oraz autonomicznej jeździe. Spółką zależną od Huawei jest **HiSilicon** koncentrujący się na tworzeniu układów scalonych do przetwarzania AI.

SenseTime to chińska firma specjalizująca się w **komputerowym przetwarzaniu obrazu**. Jej rozwiązania AI są wykorzystywane w monitoringu wideo, analizie obrazu, rozpoznawaniu twarzy, autonomicznej jeździe i sektorze finansowym. Firma rozwija również własny duży model językowo-wizualny – **SenseNova**, przeznaczony do generatywnego AI (w tym obrazy, tekst, kod). SenseTime oferuje platformę deweloperską i chmurową umożliwiającą wdrażanie modeli w różnych branżach.

DeepSeek rozwija własną rodzinę otwartych modeli językowych o nazwie **DeepSeek-VL (Vision-Language)** oraz **DeepSeek-Coder** – model dedykowany do generowania i rozumienia kodu programistycznego. Modele te charakteryzują się wysoką wydajnością przy relatywnie niskich kosztach trenowania. W testach benchmarkowych DeepSeek osiąga porównywalne wyniki do GPT-4 i Claude 2, co czyni go realnym graczem w wyścigu o dominację w obszarze open-source AI.

DeepSeek korzysta z chińskiej infrastruktury chmurowej (najprawdopodobniej wspieranej przez Alibaba Cloud lub Huawei Cloud), ale samodzielnie projektuje i trenuje modele w języku chińskim i angielskim. Firma stawia na **otwartość** (ang. **open source**) – udostępnia modele i ich dokumentację społeczności naukowej i deweloperskiej, co czyni ją chińskim odpowiednikiem podejścia Meta (**LLaMA**) lub francuskiego modelu językowego **Mistral**.

Analiza łańcucha wartości w Chinach jest dużo trudniejsza i obarczona większym ryzykiem błędów niż analizy dokonywane w innych krajach, ze względu na specyfikę ujawniania informacji w tym kraju.

W kontekście Chin analitycy [Goldman Sachs](#) prognozują, że AI może zwiększyć roczne tempo wzrostu PKB o 0,2–0,3 pp do 2030 roku. Chińskie firmy rozwijają własne modele przy niższych kosztach, co może przyspieszyć wdrażanie AI. Jednak struktura gospodarki – oparta w większym stopniu na przemyśle i

rolnictwie – ogranicza skalę korzyści w porównaniu z gospodarkami usługowymi. Adopcja AI może też znacząco wpłynąć na rynki finansowe, zwiększając wartość chińskich akcji i indeksów giełdowych (oto [link](#) do kolejnego raportu Goldman Sachs).

Na poziomie dokumentów strategicznych widoczna jest integracja AI z planem Made in China 2025 i strategią technologicznej samowystarczalności. Poniżej kilka istotnych dokumentów opisujących elementy strategii AI w Chinach:

- Angielska [analiza strategii](#) AI i tłumaczenie,
- [Raport US Chamber of Commerce](#),
- [Raport](#) Centre for Security and Emerging Eechnology,
- Według [danych Tsinghua University](#) (2023), w ciągu dekady kraj ten stał się liderem liczby publikacji naukowych z AI.

Chiny i USA łącznie odpowiadają za lwią część wszystkich światowych patentów związanych z AI, przy czym tempo przyrostu chińskich zgłoszeń patentowych rośnie stabilnie i istotnie, a wzrost ten utrzymuje się na istotnym poziomie.

2.6. Globalne wnioski o charakterze uniwersalnym

Istnieje grupa spostrzeżeń, które są istotne z punktu widzenia członków rad nadzorczych, jako istotnie wpływających na otoczenie konkurencyjne firm. Nie są one zależne od obszaru geograficznego ani wykorzystywanej technologii. Poniższa lista ma charakter niewyczerpujący i poglądowy.:

- Sztuczna inteligencja jest traktowana w wielu krajach jako narodowy priorytet gospodarczy (w tym w USA i w Chinach), a prędkość jej rozwoju znacznie przekracza tempo uchwalania regulacji.
- Sztuczna inteligencja to technologia, która licząc od momentu jej upowszechnienia wywiera największy wpływ na społeczeństwa. Jest to w szczególności widoczne w obszarze generowania treści tekstowych, obrazów i materiałów wideo.
- W każdym kraju na świecie jakość wyników modeli AI wciąż się poprawia. Przykłady to lepszy kod generowanych programów, lepszej jakości materiały tekstowe oraz video.
- Technologie AI są coraz powszechniej spotykane w otaczających nas przedmiotach – począwszy od telefonów, urządzeń medycznych, a skończywszy na autonomicznych taksówkach.
- Gigantyczne środki finansowe są inwestowane w rozwój sztucznej inteligencji w USA i Chinach, kwoty te znacznie przekraczają wydatki w Europie.
- Choć modele AI z USA są wciąż uważane za najlepsze i najwydajniejsze, różnica w stosunku do modeli z Chin szybko zmniejsza się.
- Koncepcja zaufanej sztucznej inteligencji to tzw. Responsible AI (RAI). Ta koncepcja rozwija się nierównomiernie. Wzrasta liczba odnotowanych incydentów nieprawidłowego działania AI, lecz koncepcje RAI nie są zbyt często uwzględniane przez twórców rozwiązań. Organizacje takie jak OECD, UE, ONZ a także Unia Afrykańska intensyfikują swoje działania w tym zakresie, ale uwzględnianie tych zasad przez przedsiębiorstwa nie jest satysfakcjonujące.
- Wzrasta wydajność modeli, zatem koszt użytkowania modeli i zadawania poszczególnych zapytań spada, sprawiając, że ich dostępność cenowa i penetracja rynku wzrasta.
- Rządy wielu krajów, w tym poza UE podejmują coraz więcej działań, aby uregulować obszar AI.
- Sektor edukacji boryka się z rosnącymi problemami, jeśli chodzi o gotowość nauczycieli do szkolenia, co przekłada się na dostępność wykwalifikowanych kadr.

- W chwili obecnej większość ważnych modeli AI została wytworzona przez graczy biznesowych a nie przez środowisko akademickie. Wynika to z dojrzewania rynku. (inwestorzy finansują projekty komercyjne a nie badawcze)
- Istotność AI została podkreślona przez przyznanie Nagród Nobla w obszarach fizyki i chemii. Demis Hassabis i John Jumper z Google DeepMind otrzymali Nagrodę Nobla w dziedzinie chemii za pionierską pracę nad fałdowaniem białek przy użyciu systemu AlphaFold. Tymczasem John Hopfield i Geoffrey Hinton zostali uhonorowani Nagrodą Nobla w dziedzinie fizyki za fundamentalny wkład w rozwój sieci neuronowych.
- AI potrafi rozwiązywać skomplikowane problemy matematyczne na poziomie międzynarodowych olimpiad matematycznych, ale wciąż popełnia proste z punktu widzenia ludzi błędy w innych dziedzinach.

Obszerny raport, z którego zaczerpnięto powyższe informacje to [Stanford HAI, 2024 AI Index Report 2024](#).

2.7. Możliwa odpowiedź firm europejskich

Pomimo zaległości i dominacji USA i Chin, europejskie firmy mają potencjał, by skutecznie konkurować w wybranych obszarach AI, wykorzystując unikalne atuty Starego Kontynentu.

Strategiczne inwestycje i konsolidacja zasobów

Obecnie UE inwestuje około 1/3 tego co USA w prywatny sektor AI (OECD, 2023). [Raport Europejskiej Rady ds. Innowacji](#) (2023) wskazuje, że europejskie firmy powinny skoncentrować się na konsolidacji zasobów. Europejskie inwestycje w AI są bardzo rozdrobnione. Skutkuje to duplikacją wysiłków i nieefektywnym wykorzystaniem zasobów. Firmy europejskie mogą osiągnąć znacznie więcej poprzez „strategic pooling” – łączenie zasobów w kluczowych projektach badawczych i infrastrukturalnych. Warto również zauważyć, że większość inwestycji w USA to inwestycje prywatne, ze względu na poziom rozwoju rynku kapitałowego i dostępność inwestycji PE/VC. Europa ma do zaoferowania praktycznie głównie duże inwestycje unijne. Mają one nieco inną charakterystykę niż środki inwestorów prywatnych.

Przykładem udanej inicjatywy jest projekt GAIA-X, który dąży do stworzenia europejskiej infrastruktury chmurowej umożliwiającej bezpieczne przetwarzanie danych i rozwój AI.

Współpraca nauki, biznesu i instytucji publicznych

Tworzenie konsorcjów AI, takich jak „AI4EU” czy ELLIS (European Laboratory for Learning and Intelligent Systems), ma duży potencjał. Integracja środowisk naukowych, biznesowych i administracji publicznej może przyspieszyć transfer wiedzy i technologii. Na razie inicjatywy te wyglądają bardzo ciekawie, ale są mało znane i mają nieduże budżety.

Rozwijanie talentów

Europa posiada świetną bazę akademicką i badawczą, ale przegrywa w utrzymaniu talentów. Obserwując sukcesy w tworzeniu miejsc pracy dzięki komercjalizacji pracy akademickiej, kluczowe jest stworzenie silniejszych połączeń między uczelniami a przemysłem oraz programów zatrzymujących najlepszych specjalistów AI.

Pozytywnym przykładem jest francuska inicjatywa AI for Humanity, która przyciągnęła międzynarodowe inwestycje w badania nad AI, w tym otwarcie laboratoriów badawczych przez Google, Meta i Samsung.

To powiedziawszy, warto zapoznać się z wnioskami z raportu [AI in the EU: 2024 Trends and Insights from LinkedIn](#)

Według tego raportu, tylko 0,41% pracowników w UE posiada umiejętności związane z AI, co jednak stanowi wzrost o 126% w porównaniu do 2016 roku. Dla porównania, w Wielkiej Brytanii odsetek ten wynosi 0,35%, a w USA 0,34%. Aż 42% Europejczyków nie posiada podstawowych umiejętności cyfrowych, co stanowi poważną barierę dla rozwoju AI.

Raport podkreśla, że mimo wzrostu liczby specjalistów AI, UE nadal boryka się z poważnymi wyzwaniami. Jest nim zwłaszcza niski odsetek pracowników z umiejętnościami AI ograniczający możliwości rozwoju technologii. Istotną kwestią są również braki w edukacji cyfrowej – duża część społeczeństwa nie posiada podstawowych umiejętności cyfrowych, co utrudnia adaptację nowych technologii.

Aby sprostać tym wyzwaniom, raport zaleca inwestycje w edukację i szkolenia: rozwijanie programów edukacyjnych z zakresu AI i umiejętności cyfrowych oraz łączenie działań rządów, przemysłu i instytucji edukacyjnych w celu tworzenia sprzyjającego ekosystemu dla AI.

Wykorzystanie unikalnych europejskich atutów

Europa posiada kilka unikalnych przewag, które mogą stanowić podstawę konkurencyjnej strategii:

Silna baza przemysłowa

Jak wskazuje [raport Boston Consulting Group](#) (2023), europejskie firmy mogą wykorzystać swoją silną pozycję w przemyśle: "Europejskie doświadczenia w sektorach takich jak motoryzacja, chemia i produkcja zaawansowana, stwarzają unikalną możliwość rozwoju przemysłowej AI (Industrial AI), gdzie Europa może osiągnąć globalną przewagę konkurencyjną."

Standardy etyczne i prawne

GDPR (RODO) i AI Act (2024) mogą stać się światowym wzorcem odpowiedzialnego rozwoju AI.

"Europejskie podejście do AI koncentrujące się na etyce, przejrzystości i odpowiedzialności, może stać się globalnym standardem dla odpowiedzialnej AI, szczególnie w sektorach wrażliwych, takich jak opieka zdrowotna czy finanse" (Komisja Europejska, 2023). Różnorodność kulturowa – AI trenowana na zróżnicowanych danych – może zapewnić większą inkluzywność i mniejszą podatność na uprzedzenia.

Kształtowanie regulacji AI

Europa powinna dalej rozwijać swoją wizję „AI godnej zaufania” (Trusted AI), jak definiuje ją Komisja Europejska. AI Act to ważny krok, ale nie koniec drogi. Firmy europejskie powinny aktywnie uczestniczyć w kształtowaniu tych regulacji, aby zapewnić równowagę między ochroną wartości, a wspieraniem innowacji.

Specjalizacja w niszach o strategicznym znaczeniu

Raport Europejskiego Instytutu Innowacji i Technologii (2023) sugeruje:

"Europejskie firmy powinny skupić się na obszarach, w których mogą uzyskać przewagę komparatywną, takich jak AI w ochronie zdrowia, zrównoważonym rozwoju i przemyśle, zamiast bezpośrednio konkurować z gigantami technologicznymi w dziedzinie AI konsumenckiej."

Przykłady specjalizacji:

- AI w medycynie (np. Siemens Healthineers)
- Automatyzacja przemysłu (np. ABB)
- AI dla zrównoważonego rozwoju (np. aplikacje do zarządzania i oszczędzania energii w budynkach)

Przykładem europejskiego sukcesu jest założona przez Polaka niemiecka firma DeepL, która stworzyła jeden z najlepszych systemów tłumaczenia maszynowego, konkurując skutecznie z rozwiązaniami Google'a czy Microsoftu

Tworzenie europejskich hubów AI

Inicjatywy takie jak „GAIA-X”, „European AI Startups Landscape” czy „European AI Valleys” pokazują kierunek tworzenia regionalnych centrów kompetencji, które mogą przyciągać talenty i inwestycje z całego świata.

2.8. Bezpieczeństwo geopolityczne

Rywalizacja w dziedzinie AI wykracza poza biznes – to kwestia **bezpieczeństwa narodowego i suwerenności technologicznej**. USA i Chiny traktują AI jako strategiczną infrastrukturę, a Europa musi wypracować własną odpowiedź. Według [analiz NATO Strategic Communications Centre of Excellence](#) (2023), zależność od zagranicznych technologii AI może stwarzać ryzyko dla bezpieczeństwa narodowego państw europejskich:

"Zdolność Europy do samodzielnego działania w zakresie sztucznej inteligencji to nie tylko kwestia konkurencyjności, lecz przede wszystkim suwerenności i możliwości kształtowania naszej własnej przyszłości."

„Europe's ability to act independently in AI is not just about competitiveness, but about sovereignty and the capacity to shape our own future."

– Thierry Breton, Komisarz UE ds. Rynku Wewnętrznego

Europa musi uniezależnić się od zewnętrznych dostawców AI, zwłaszcza w zakresie infrastruktury chmurowej, półprzewodników i dużych modeli językowych (LLM).

2.9. Sytuacja Polski i Europy Środkowo-Wschodniej

Kraje naszego regionu posiadają duże zasoby zmotywowanego i wykształconego personelu IT. Polska charakteryzuje się wysoką penetracją technologii cyfrowych, takich jak: dobra sieć telekomunikacyjna, wysoka penetracja e-commerce. Kraje naszego regionu, a zwłaszcza Polska, stworzyły wiele unikalnych rozwiązań cyfrowych, zarówno w sferze publicznej, jak i prywatnej (kilka przykładów w Polsce to: mObywatel, istotna pozycja rynkowa Allegro, znakomita bankowość internetowa, system płatności BLIK, a ostatnio sukces modelu językowego Bielik. Najnowszym sukcesem jest pozyskanie finansowania i ekspansja zagraniczna polskiej firmy ElevenLabs. Przykłady innych firm w regionie to UiPath w Rumunii czy Revolut, który działa z UK, ale ma wschodnioeuropejskie korzenie).

Niestety, nie mamy jednak wystarczających zasobów kapitałowych, aby inwestować w centra danych, czy dostateczny rozwój własnych technologii. Polska (a tym bardziej inne kraje regionu) jest również za małym krajem, aby zgromadzić wielkie ilości danych, które mogłyby stanowić przewagę konkurencyjną dla budowy globalnych rozwiązań AI. Rynek kapitałowy jest również zbyt słaby, aby takie projekty samodzielnie sfinansować.

Siłą rzeczy tworzy to zależność od zagranicznych dostawców kluczowych technologii, pozostawiając nam rodzaj kreatywnej integracji. Polskie firmy informatyczne koncentrują się zatem na bardziej niszowych projektach dotyczących wybranych obszarów technologii. W czasach odwrotu od globalizacji i kwestionowania modelu wolnego handlu, gospodarka naszego kraju i regionu powinna skupić się szczególnie na wykorzystaniu i wspieraniu takich rodzimych technologii, aby rozwijać kompetencje i niezależność technologiczną tam, gdzie jest to tylko możliwe.

2.10. Podsumowanie sytuacji geopolitycznej

Zrozumienie aktualnej pozycji Europy w kontekście globalnej rywalizacji AI jest niezbędne dla członków rad nadzorczych. Raport Draghiego stanowi jasny sygnał alarmowy dla Europy w kontekście rozwoju AI. Choć dominacja USA i Chin jest niezaprzeczalna, europejskie firmy mają szansę na konkurowanie poprzez:

- Strategiczne połączenie zasobów i konsolidację inwestycji
- Wykorzystanie unikalnych europejskich atutów, jak silna baza przemysłowa i podejście oparte na wartościach
- Budowanie ekosystemu talentów łączącego naukę z biznesem
- Specjalizację w strategicznych niszach, gdzie Europa może uzyskać przewagę konkurencyjną.

3. Jak przygotować się do nadzoru nad AI?

Efektywne sprawowanie nadzoru nad sztuczną inteligencją (AI) wymaga konkretnej wiedzy i praktycznych kroków. Członkowie i członkinie rad nadzorczych powinni zrozumieć specyfikę technologii, aktualne trendy oraz znaczenie danych, które są kluczowe dla wdrażania skutecznych rozwiązań AI. W zależności od branży przedsiębiorstwa i roli w radzie nadzorczej może to być poziom bardziej ogólny, zapewniający zrozumienie kluczowych wyzwań lub zaawansowany, jeśli spółka działa w branży ze szczególną ekspozycją na AI, zarówno w zakresie szans jak i ryzyk.

3.1. Podstawowe kwestie AI w spółkach – co trzeba wiedzieć?

Według badań Boston Consulting Group 3 z 4 członków kadry menedżerskiej klasyfikuje AI jako jeden z trzech najwyższych priorytetów strategicznych. Jednocześnie tylko 1 z 4 z nich potrafi wskazać konkretną wartość, którą buduje w firmie AI.

Ugruntowane i sprawdzone sposoby budowania wartości przedsiębiorstw za pomocą AI to:

- wdrażanie sztucznej inteligencji w codziennych zadaniach, aby osiągnąć 10–20% wzrostu wydajności;
- modyfikacja kluczowych procesów, aby poprawić efektywność i skuteczność o 30–50%;
- tworzenie nowych produktów i usług, aby zbudować długoterminową przewagę konkurencyjną.

Pobieżna analiza ryzyk i korzyści wskazuje, że największą korzyść można odnieść w obszarze drugim, gdzie istnieje wysoki potencjał zwrotu z inwestycji przy niskim ryzyku związanym ze znajomością materii.

Wiele firm stosuje strategię „spray and pray” pozwalając na wprowadzenie wielu niewielkich pilotaży różnych narzędzi AI (akceleracje, sandboxy, konkursy), nie koncentrując wysiłków i nie dając im należycie wysokiego priorytetu w organizacji. Istnieje wysokie ryzyko, że takie podejście nie przyniesie najlepszych możliwych rezultatów. Wiąże się to także po części z brakiem czytelnej artykulacji oczekiwań i mierzalnych celów przy wdrażaniu rozwiązań AI

Istnieją 4 główne domeny ryzyk związanych z projektami AI:

1. ryzyko prawne,
2. ryzyko technologiczne,
3. ryzyko organizacyjne,
4. ryzyko społeczne i etyczne.

Organizacje dostrzegają ryzyka związane z odpowiedzialnym wykorzystaniem AI (RAI), ale działania zaradcze często są opóźnione. Według badania McKinsey, choć wiele firm rozpoznaje kluczowe zagrożenia, takie jak niedokładność, kwestię zgodności z przepisami czy cyberbezpieczeństwo, tylko odpowiednio 64%, 63% i 60% respondentów przyznaje, że stanowią one dla nich realne obawy – w efekcie nie wszystkie firmy podejmują konkretne kroki w celu ich ograniczenia.

Na całym świecie w 2024 roku zauważalnie wzrosło zainteresowanie tworzeniem zasad odpowiedzialnego zarządzania AI. Międzynarodowe instytucje, takie jak OECD, Unia Europejska, ONZ czy Unia Afrykańska, opublikowały ramy regulacyjne dotyczące przejrzystości, wytłumaczalności oraz wiarygodności systemów AI.

Jednocześnie gwałtownie kurczy się wspólna przestrzeń danych. Modele AI potrzebują ogromnych ilości publicznie dostępnych danych z internetu, ale od 2023 do 2024 roku nastąpił wyraźny wzrost ograniczeń w zakresie ich wykorzystania. W popularnym zbiorze danych C4 udział „zablokowanych” treści wzrósł z 5–7% do 20–33%, co może wpłynąć na różnorodność danych, możliwości dopasowania modeli oraz ich skalowalność – a także wymusić nowe podejścia do nauki przy ograniczonych zasobach.

Pomimo prób eliminowania stronniczości, zaawansowane modele językowe (LLM), takie jak GPT-4 czy Claude 3 Sonnet, nadal wykazują ukryte stronniczość i subiektywność (**bias**). Modele te np. częściej kojarzą osoby czarnoskóre z negatywnymi określeniami, przypisują kobietom role w naukach humanistycznych a nie w ścisłych, a mężczyzn faworyzują w kontekście przywództwa – tym samym wzmacniając stereotypy m.in. rasowe i płciowe. Mimo że wskaźniki uprzedzeń na standardowych benchmarkach uległy poprawie, problem stronniczości w AI pozostaje powszechny.

3.2. Jakie trendy przeważają w AI w 2025 roku?

AI to bardzo dynamicznie zmieniający się obszar, dlatego rady nadzorcze muszą śledzić na bieżąco najważniejsze trendy. Kluczowe trendy w 2025 roku, które rady nadzorcze powinny uwzględnić w swoich planach strategicznych to:

Regulacje prawne

W 2025 roku organizacje muszą przygotować się na coraz szerszy zakres regulacji dotyczących sztucznej inteligencji.

Europejski AI Act

- Wprowadza klasyfikację systemów AI według poziomu ryzyka: minimalne, ograniczone, wysokie i niedopuszczalne.
- Branże objęte: finanse, edukacja, ochrona zdrowia, zatrudnienie, transport, usługi publiczne. W przypadku systemów wysokiego ryzyka (np. scoring kredytowy, analiza kandydatów do pracy, systemy medyczne), firmy muszą spełniać surowe wymogi dot. jakości danych, przejrzystości, nadzoru ludzkiego.
- Źródła do monitorowania: AI Act Tracker (Future of Life Institute), kanał YouTube Ministerstwa Cyfryzacji.

Dyrektywa NIS2

- Dotyczy cyberbezpieczeństwa podmiotów kluczowych dla funkcjonowania państwa – m.in. energetyka, transport, finanse, zdrowie, dostawcy infrastruktury cyfrowej.
- Nakłada obowiązki w zakresie zarządzania ryzykiem IT i zgłaszania incydentów.

- **Ważne dla rad nadzorczych:** firmy powinny zidentyfikować, czy podlegają dyrektywie i dostosować systemy bezpieczeństwa AI.

Dyrektywa DORA (Digital Operational Resilience Act)

- Skierowana do instytucji finansowych i fintechów.
- Zobowiązuje do stosowania odpornych na zakłócenia rozwiązań IT, również w kontekście wykorzystywanych modeli AI.
- Wymaga audytowalności, ciągłości działania i testów odporności cyfrowej.

ISO/IEC 23894:2023

- Nowy standard dotyczący zarządzania ryzykiem AI.
- Zalecany dla firm wdrażających lub rozwijających własne systemy AI – zwłaszcza w sektorach regulowanych.

AI Explainability (wyjaśnialność)

Explainability to zdolność systemu AI do wyjaśnienia, jak i dlaczego doszedł do konkretnej decyzji.

Dlaczego to ważne?

- Wysoka wyjaśnialność zwiększa zaufanie interesariuszy, zgodność z regulacjami i zmniejsza ryzyko reputacyjne.
- Pozwala uniknąć sytuacji, w których firma nie potrafi wytłumaczyć klientowi decyzji np. o odmowie kredytu lub rekomendacji leczenia.

Przykłady

- Bank odrzucający wniosek kredytowy na podstawie „czarnej skrzynki” (black-box model), bez możliwości wskazania konkretnego powodu (sytuacja analizowana przez Fundację Panoptikon).
- Firma rekrutacyjna korzystająca z AI do preselekcji kandydatów – w przypadku braku **explainability**, może dojść do niesprawiedliwego odrzucenia kandydatów z określonych grup.

Cyberbezpieczeństwo AI

Nowe generacje zagrożeń wynikają ze zdolności cyberprzestępców do wykorzystania AI.

Dlaczego teraz bardziej niż kiedykolwiek?

- AI obniża próg wejścia dla ataków – automatyzacja phishingu, deepfake, spoofing, generowanie złośliwego kodu.
- Pracownicy często nieświadomie udostępniają dane w LLM-ach (np. ChatGPT), co może skutkować ich wyciekiem.

Przykłady

- Wyciek dokumentów korporacyjnych poprzez pracownika testującego GenAI.

- Podszywanie się pod dział wsparcia klienta – ataki socjotechniczne wygenerowane przez AI.

Na co zwrócić uwagę w nadzorze?

- Czy firma, która obsługuje spółkę w zakresie cyberbezpieczeństwa, prowadzi aktualne szkolenia dotyczące zagrożeń AI?
- Czy są wdrożone zasady korzystania z narzędzi AI przez pracowników (np. zakaz wklejania danych osobowych do LLM)?

Etyka AI

Dlaczego to istotne?

- Algorytmy mogą wzmacniać istniejące uprzedzenia powodując ryzyko niezamierzonej dyskryminacji (np. względem płci, wieku, rasy).
- AI może naruszać prywatność – np. zbierając zbyt szeroki zakres danych bez zgody.

Przykłady

- **Amazon** zrezygnował z AI w rekrutacji z uwagi na faworyzowanie mężczyzn.
- **Allegro** wdrożyło procedury przeglądu algorytmów rekomendacji.
- **mBank** monitoruje scoring kredytowy pod kątem zgodności z zasadami etycznymi.

Rekomendacje dla rad nadzorczych

- Wyznaczyć osobę odpowiedzialną za regularne raportowanie o zmianach prawnych (może być to członek działu compliance).
- Audyty bezpieczeństwa systemów AI przez wyspecjalizowane firmy.
- Organizowanie cyklicznych warsztatów z zakresu etyki AI (np. firma zewnętrzna lub uniwersytety).
- Monitorowanie zmian prawnych i nowych interpretacji np. na stronie Ministerstwa Cyfryzacji lub poprzez newslettery np. dużych firm doradczych – sekcje **Legal** i **Regulatory**.

3.3. GenAI vs Data AI – znaczenie danych

Rodzaje AI

Ostatnie dwa lata to czas ogromnej popularności tzw. Generatywnej AI zapoczątkowanej przez udostępnienie Chata-GPT przez OpenAI. Gwałtownie rośnie rola tzw. **AI Agents**. Jednak wciąż ogromna większość skutecznych i opłacalnych wdrożeń AI w Polsce i na świecie to właśnie projekty Data AI. Stąd warto mieć świadomość tego rozróżnienia i właściwie balansować zainteresowanie najnowszymi produktami typu Chat-GPT względem realnych możliwości, jakie daje skupienie na analizie danych (określane czasami jako Tradycyjne AI albo Data AI)

Generatywna AI (GenAI)

Tworzy nowe treści (np. teksty, obrazy, dźwięki). Przykłady: ChatGPT (generowanie tekstów), Midjourney (generowanie obrazów), polska aplikacja Lekta (syntezator mowy). Najbardziej rozpoznawalne są obecnie tzw. duże modele językowe (ang. LLM – large language models) – przykłady to Chat-GPT, Llama, Gemini, DeepSeek. GenAI wykorzystywana jest m.in. w marketingu, obsłudze klienta, HR czy tworzeniu materiałów szkoleniowych. Jej możliwości są imponujące, ale jej efektywne wdrożenie wymaga zrozumienia ograniczeń, takich jak halucynacje, nieprzewidywalność czy brak pełnej kontroli nad wygenerowaną treścią.

Data AI

Analiza danych, zarówno liczbowych jak i np. wizualnych, do podejmowania decyzji biznesowych. Przykłady zastosowań to:

- modele predykcyjne, np. przewidywanie odejścia klientów w bankowości lub telekomunikacji,
- systemy rekomendacji produktów, jak w Allegro, Amazonie czy Netflixie,
- optymalizacja łańcuchów dostaw, np. w sieciach takich jak Żabka czy CCC,
- Analiza obrazu i sygnałów, np. do diagnostyki medycznej lub monitoringu produkcji.

AI Agents

Rośnie znaczenie **AI Agents**, czyli inteligentnych agentów, które nie tylko przetwarzają dane lub generują treści, ale aktywnie działają w imieniu użytkownika, podejmując decyzje i wykonując zadania w systemach IT. Agenci AI potrafią m.in. zaplanować i wykonać ciąg działań (np. zarezerwować podróż służbową, przeanalizować dane finansowe i wygenerować raport, zintegrować się z kalendarzem, CRM czy bazą maili). Technologie te często łączą GenAI z Data AI oraz z komponentami automatyzacji procesów – co stanowi krok w kierunku tzw. autonomicznych systemów wspomagania decyzji.

Inne zastosowania

AI pozwala na różne inne zastosowania, np. prowadzenie procesów biznesowych przez roboty – tzw. RPA (**Robotic Process Automation**), które umożliwia zautomatyzowanie powtarzalnych procesów takich jak księgowanie faktur czy weryfikacja przelewów.

Przykłady zastosowań w Polsce

Data AI

- Banki takie jak PKO BP: analiza danych do personalizacji ofert kredytowych.
- Portale e-commerce takie jak Allegro: algorytmy rekomendacyjne zwiększające sprzedaż.
- Firmy ubezpieczeniowe takie jak PZU: Szacowanie kosztów szkód na podstawie zdjęć samochodów.

GenAI

- Firmy usługowe takie jak InPost: chatbot do obsługi klientów wykorzystujący GenAI (automatyzacja komunikacji).
- Santander Bank Polska: tworzenie spersonalizowanych wiadomości marketingowych przy użyciu GenAI.
- Wsparcie zespołów programistycznych w generowaniu kodu przy użyciu LLM.

Rekomendacje dla rad nadzorczych

Weryfikacja gotowości spółki pod kątem jakości i dostępności danych do prowadzenia analiz oraz realizacji decyzji na ich podstawie. Obejmuje to sprawdzenie czy systemy danych są ze sobą odpowiednio zintegrowane oraz czy podejmowane decyzje biznesowe mogą być efektywnie wdrażane w konkretne działania (np. możliwość technologiczna dynamicznego kształtowania cen).

Zatrudnienie lub wyznaczenie Chief Data Officer (CDO) odpowiedzialnego za zarządzanie danymi i zapewnienie ich jakości, szczególnie jeśli spółka decyduje się na intensywną pracę z danymi.

Prowadzenie regularnych rozmów z przedstawicielami innych spółek oraz dostawcami technologii (np. Google, Microsoft, NVIDIA, Salesforce), dotyczących aktualnych wdrożeń sztucznej inteligencji, ich rezultatów oraz głównych wyzwań związanych z tymi projektami.

3.4. Skąd się uczyć o AI? Szkolenie jako całościowa odpowiedzialność rad nadzorczych i kadry

Zrozumienie AI to nie tylko obowiązek działów IT, lecz także strategiczna odpowiedzialność rady nadzorczej. Członkowie rady mogą nie mieć specjalistycznej wiedzy technicznej, dlatego kluczowe jest zapewnienie sobie praktycznych i przystępnych źródeł informacji i regularne sięganie do nich, ponieważ zmiany i postęp są tu szczególnie dynamiczne.

Rekomendacje dla rad nadzorczych

Samodzielna edukacja

Poniższe treści to przykładowe propozycje, które warto rozważyć. Możliwości jest jednak znacznie więcej i zachęcamy do uzupełnienia tej listy własnymi wyborami.

- Kursy online: „[AI for Everyone](#)” (Coursera, Andrew Ng), dostępne także po polsku.
- Podcasty: „Biznes Myśli” (po polsku), „Technologicznie Podcast” Jarosława Kuźniara (po polsku), „NVIDIA AI Podcast” (po angielsku).
- Książki: Ethan Mollick „Co-Intelligence”, Max Tegmark „Życie 3.0” (dostępna po polsku).

Regularne prezentacje ekspertów wewnętrznych

- Organizowane przez przewodniczącego rady, członków rady lub kierownictwo spółki spotkania kwartalne (np. 1-godzinne sesje) z udziałem CTO, CIO, Dyrektora ds. Ryzyka lub osoby odpowiedzialnej za cyfryzację.

- Forma: krótkie prezentacje (do 20 minut) + Q&A (30-40 minut), skupiające się na praktycznych przykładach wdrożeń AI w firmie i ich efektach biznesowych.

Korzystanie z ekspertów zewnętrznych

- Kontakt z firmami doradczymi w celu zapytania o możliwość przeprowadzenia warsztatów lub indywidualnych konsultacji dotyczących wdrożenia AI.
- Firmy te są zainteresowane współpracą ze względu na możliwe długoterminowe relacje biznesowe oraz budowanie renomy poprzez współpracę z dużymi spółkami. Często organizują też większe wydarzenia w tematyce nowych technologii, w tym AI, na które mogą zaprosić członków rady nadzorczej.
- Współpraca z polskimi uczelniami technicznymi (np. Politechnika Warszawska, Politechnika Wrocławska), które prowadzą specjalistyczne warsztaty AI dla biznesu.

Odpowiedzialność i zaangażowanie

- Przewodniczący rady, wyznaczony członek rady lub zarządu może stworzyć prosty plan edukacyjny, obejmujący konkretne terminy i tematy.
- Regularne uwzględnianie tematu AI w planie posiedzeń rady nadzorczej pomoże utrzymać dynamikę edukacji i dyskusji.

4. Fundamenty efektywnego nadzoru nad AI w organizacji

4.1. Istotność strategiczna i ryzyko braku działania

4.1.1. Bilans otwarcia AI w organizacji

Punkt wyjścia dla członków rad nadzorczych to rzetelna ocena aktualnej pozycji organizacji w zakresie wykorzystania AI oraz rozumienie konsekwencji strategicznych, operacyjnych i konkurencyjnych braku działania.

Kluczowe pytania

- Jakie inicjatywy AI są obecnie realizowane w organizacji i jaki dają efekt biznesowy?
- Jaka jest nasza docelowa pozycja w zakresie wykorzystania AI w perspektywie krótko- i długoterminowej?
- Jakie konkretne wskaźniki mierzą nasz postęp w obszarze AI?
- Jakie są potencjalne konsekwencje strategiczne, operacyjne i konkurencyjne zaniechania inwestycji w AI?
- Jakie trendy rynkowe i działania konkurencji w obszarze AI mogą wpłynąć na naszą pozycję?

Potencjalne działania

- Przeprowadzenie audytu obecnego stanu zaawansowania AI w organizacji.
- Identyfikacja kluczowych obszarów potencjalnego zastosowania AI oraz określenie aspiracji strategicznych w tym obszarze.
- Ocena ryzyka strategicznego związanego z brakiem aktywnego podejścia do AI, w tym potencjalnej utraty przewagi konkurencyjnej.



Przykład pytania

Jakie cele strategiczne możemy osiągnąć dzięki AI i jakie ryzyka wynikają z opóźnienia w jej wdrożeniu?



Przykład działania

Przeprowadzenie analizy benchmarkowej z konkurencją w celu oceny pozycji firmy oraz stworzenie mapy drogowej implementacji AI.

4.1.2. Umiejscowienie odpowiedzialności

Sprawowanie efektywnego nadzoru wymaga jasnego określenia odpowiedzialności za AI na wszystkich poziomach organizacji.

Kluczowe pytania

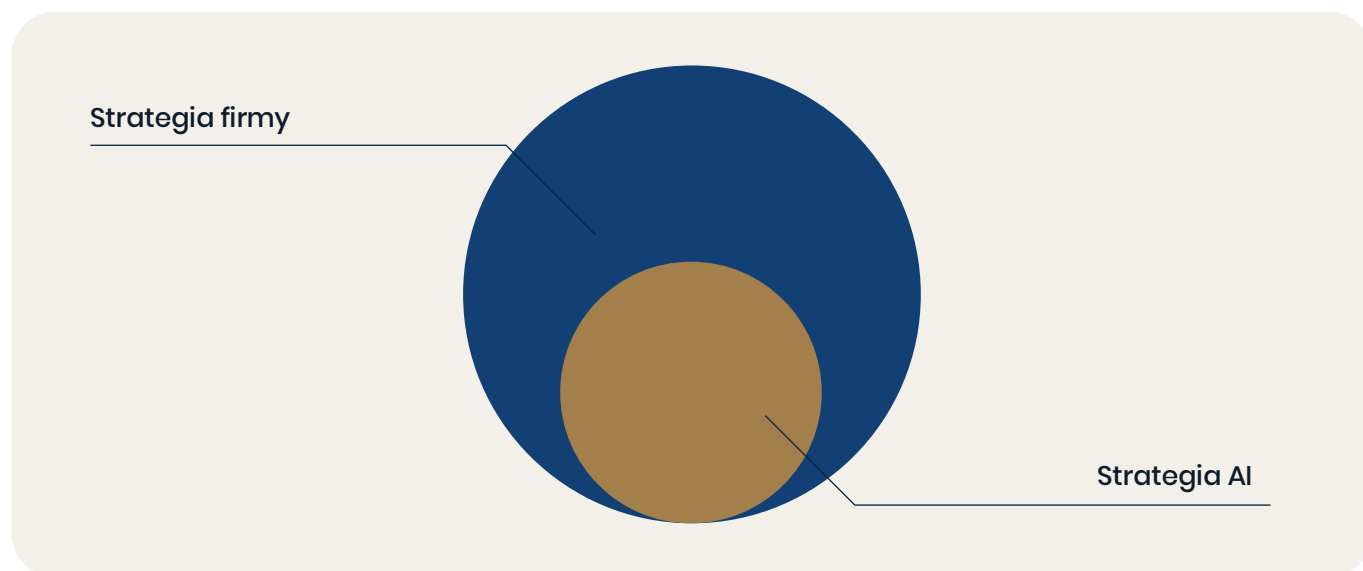
- Kto na poziomie zarządu jest ostatecznie odpowiedzialny za strategię AI i jej wdrożenie i czy cała rada nadzorcza sprawuje nad tym nadzór czy też wspiera ją w tym któryś z jej komitetów?
- Czy istnieją dedykowane komitety lub role odpowiedzialne za AI?
- Czy odpowiedzialność za AI jest traktowana jako element tymczasowy czy stały w strukturze zarządzania?

Możliwe działania

- Formalne określenie i udokumentowanie odpowiedzialności za AI na wszystkich szczeblach zarządzania.
- Rozważenie powołania dedykowanego zespołu lub komitetu ds. AI na poziomie rady nadzorczej.
- Ustanowienie regularnych raportów z obszaru AI jako stałego elementu posiedzeń rady nadzorczej.

4.1.3. Osadzenie AI w strategii organizacji

AI nie powinna funkcjonować jako izolowany projekt, lecz jako element zintegrowany z całościową strategią biznesową.



Kluczowe pytania

- W jaki sposób AI przyczynia się do realizacji kluczowych celów strategicznych?
- Czy strategia AI jest spójna z innymi inicjatywami transformacji cyfrowej?
- Jak mierzymy wpływ inicjatyw AI na realizację celów strategicznych?

Potencjalne działania

- Weryfikacja spójności inicjatyw AI z ogólną strategią organizacji.
- Opracowanie wskaźników mierzących wpływ AI na realizację strategii.
- Integracja AI z cyklicznym procesem planowania strategicznego.

4.1.4. Identyfikacja możliwości biznesowych

Rada nadzorcza powinna wymagać systematycznego podejścia do identyfikacji obszarów, gdzie AI może przynieść największą wartość.

Kluczowe pytania

- Jakie są priorytety zastosowania AI w organizacji i dlaczego?
- Które procesy biznesowe mogą odnieść największe korzyści z wdrożenia AI?
- Czy istnieją nowe modele biznesowe, które możemy rozwinąć dzięki AI?

Możliwe działania

- Przeprowadzenie warsztatów identyfikujących potencjał AI w poszczególnych jednostkach biznesowych.
- Opracowanie mapy priorytetowych inicjatyw AI z oszacowaniem potencjalnych korzyści.
- Regularny przegląd nowych możliwości biznesowych związanych z ewolucją technologii AI.

4.1.5. Ekonomiczna analiza kosztów i korzyści

Inicjatywy AI wymagają rygorystycznej oceny finansowej, uwzględniającej nie tylko bezpośrednie koszty wdrożenia, ale także długoterminowy zwrot z inwestycji.

Kluczowe pytania

- Jaki jest całkowity koszt wdrożenia i utrzymania rozwiązań AI?
- Jakie są mierzalne i niemierzalne korzyści z inicjatyw AI?
- Jaki jest oczekiwany ROI i czas zwrotu z inwestycji w AI?
- Jakie ryzyka finansowe są związane z projektami AI?

Możliwe działania

- Opracowanie modelu finansowego oceniającego całkowity koszt posiadania (TCO) rozwiązań AI.
- Wdrożenie systematycznego pomiaru korzyści z wdrożonych rozwiązań AI.
- Ustanowienie progów zwrotu z inwestycji jako kryteriów akceptacji projektów AI.

4.1.6. Zarządzanie talentem i edukacja

Skuteczne wykorzystanie AI wymaga odpowiednich kompetencji na wszystkich poziomach organizacji.

Kluczowe pytania

- Czy posiadamy odpowiednie zasoby ludzkie do realizacji strategii AI?
- Jak rozwijamy kompetencje AI wśród istniejących pracowników?
- Jaka jest strategia pozyskiwania talentów w obszarze AI?

Możliwe działania

- Audyt luk kompetencyjnych w organizacji.
- Opracowanie programu szkoleń z AI dla różnych grup pracowników, w tym kadry zarządzającej i rady nadzorczej.
- Włączenie kompetencji AI do procesów rekrutacyjnych lub planów sukcesji.



Przykład pytania

Jakie kompetencje w zakresie AI są kluczowe dla naszego rozwoju?



Przykład działania

Opracowanie programów szkoleń i certyfikacji dla pracowników oraz pozyskiwanie specjalistów poprzez partnerstwa z instytucjami akademickimi.

4.1.7. Realistyczne osadzenie w możliwościach spółki

Ambicje w zakresie AI muszą być dostosowane do rzeczywistych możliwości organizacji. W obliczu przykładów zastosowań z całego świata i informacji napływających ze wszystkich stron nietrudno popaść w nadmierną presję oraz chęć dorównania najlepszym lub wdrożenia wszystkiego naraz. Rzeczywistość większości spółek nie pozwala jednak na zrównoważone wdrożenie wszystkich najlepszych praktyk w krótkim czasie i wymaga rozważnego dopasowania planu działania do możliwości organizacji.

Jeden dobrze wdrożony projekt, który przyniesie wymierne rezultaty, jest często dobrym punktem wyjścia, ponieważ uczy organizację pracy z AI, buduje zaufanie do tej technologii i staje się inspiracją do dalszych działań.

Kluczowe pytania

- Czy cele związane z AI są realistyczne w kontekście dostępnych zasobów i możliwości?
- Jaka jest dojrzałość organizacji do absorpcji rozwiązań AI?

- Czy tempo wdrażania AI jest dostosowane do możliwości adaptacyjnych organizacji?
- Co da największy efekt w możliwie krótkim czasie i w związku z tym warto od tego zacząć?

Możliwe działania

- Przeprowadzenie oceny dojrzałości organizacji w kontekście wdrażania AI.
- Ustanowienie etapowego planu rozwoju zdolności AI w organizacji.
- Regularna weryfikacja realności celów i harmonogramów związanych z AI.

4.1.8. Jakie zadania zlecamy podmiotom zewnętrznym, a co pozostaje wewnątrz organizacji (czyli build vs. buy)

Strategiczne decyzje dotyczące sourcingu kompetencji i rozwiązań AI są kluczowe dla zrównoważonego rozwoju tych technologii w organizacji. Organizacje mogą mieć swoje preferencje do tworzenia rozwiązań IT wewnątrz lub kupowania ich z zewnątrz. Mądry zespół potrafi różnicować te podejścia w zależności od potrzeby i dostępności gotowych rozwiązań na rynku. Czasami te preferencje mogą być jednak zniekształcone i zespoły mogą chętniej wybierać budowanie rozwiązań wewnątrz (np. żeby uzasadnić potrzebę istnienia i rozwijania swoich działów, zapewnić długoterminową istotność swojego działu) lub kupowanie rozwiązań na zewnątrz (np. w celu zmniejszenia odpowiedzialności, ze względu na zachęty od firm technologicznych za zakup ich rozwiązań). Kierownictwo i rada nadzorcza stoją przed trudnym zadaniem odpowiedniej oceny i doboru strategii na korzyść spółki.

W sytuacji, w której na rynku istnieje gotowe rozwiązanie obsługujące większość wymagań spółki, najlepsze rozwiązanie to zazwyczaj zakup gotowego rozwiązania. W szczególności, gdy nie dotyczy to procesów decydujących o kluczowych przewagach spółki. W sytuacji, gdy spółka wymaga specyficznej obsługi konkretnych procesów lub dany system decyduje o przewadze konkurencyjnej lub stanowi centralny aspekt działalności spółki – budowa systemu jest uzasadniona. W wypadku AI dla ogromnej większości firm nie ma sensu budowa własnego modelu językowego albo systemu CRM – wystarczy dostosowanie gotowych modeli. Dla kontrastu w wypadku firmy produkcyjnej budowa własnego modelu obsługującego istotną część procesu produkcji – np. optymalizacja proporcji składników na podstawie AI, jest zdecydowanie uzasadniona.

Kluczowe pytania

- Jakie elementy strategii AI realizujemy wewnątrz, a jakie zlecamy do realizacji podmiotom zewnętrznym?
- Jakie są kryteria podejmowania tych decyzji (koszty, kompetencje strategiczne, ryzyko)?
- Jak zarządzamy relacjami z zewnętrznymi dostawcami rozwiązań AI?

Działania

- Określenie jasnych zasad dotyczących outsourcingu i in-sourcingu w obszarze AI.
- Wdrożenie systematycznego procesu oceny dostawców i partnerów AI.

- Ustanowienie mechanizmów monitorowania efektywności współpracy z zewnętrznymi partnerami.

4.2. Odpowiedzialne AI

4.2.1. Zgodność z wartościami firmy

Wdrażane rozwiązania AI muszą odzwierciedlać fundamentalne wartości organizacji.

Kluczowe pytania

- Czy nasze inicjatywy AI są zgodne z deklarowanymi wartościami firmy?
- Jak zapewniamy, że algorytmy AI działają zgodnie z naszym kodeksem etycznym?
- Czy pracownicy rozumieją etyczne wymagania związane z wykorzystaniem AI?

Możliwe działania

- Opracowanie kodeksu etycznego AI specyficznego dla organizacji lub rozbudowanie kodeksu etyki organizacji o obszar AI.
- Włączenie kryteriów etycznych do procesu oceny inicjatyw AI.
- Wdrożenie mechanizmów zgłaszania potencjalnych naruszeń etyki AI.

4.2.2. Integracja z systemami kontroli wewnętrznej

Odpowiedzialne wykorzystanie AI wymaga włączenia odpowiednich mechanizmów kontrolnych.

Kluczowe pytania

- Czy funkcje audytu, ryzyka i compliance są zaangażowane w nadzór nad AI?
- Jakie mechanizmy kontrolne stosujemy do monitorowania systemów AI?
- Czy posiadamy procedury wykrywania i reagowania na nieprawidłowości w działaniu AI?

Możliwe działania

- Opracowanie ram kontrolnych specyficznych dla systemów AI.
- Szkolenia dla zespołów audytu i compliance w zakresie oceny ryzyk związanych z AI.
- Włączenie ryzyk AI do istniejących procesów zarządzania ryzykiem.



Przykład pytania

Jak zintegrować AI z istniejącymi systemami kontroli wewnętrznej?



Przykład działania

Utworzenie interdyscyplinarnego zespołu ds. AI, w którego składzie znajdą się przedstawiciele działów IT, prawa, audytu i zarządzania ryzykiem.

4.2.3. Polityka oceny zaufania i wiarygodności

Organizacje muszą wypracować systematyczne podejście do oceny wiarygodności stosowanych rozwiązań AI.

Kluczowe pytania

- Jakie kryteria stosujemy do oceny wiarygodności systemów AI?
- Jak weryfikujemy jakość i rzetelność wyników generowanych przez AI?
- Czy posiadamy mechanizmy audytu algorytmów pod kątem ich wiarygodności?

Możliwe działania

- Opracowanie standardów oceny wiarygodności systemów AI.
- Wdrożenie regularnych testów i walidacji rozwiązań AI.
- Ustanowienie procedur certyfikacji wewnętrznej dla krytycznych systemów AI.

4.2.4. Cyberbezpieczeństwo w kontekście AI

AI wprowadza nowe zagrożenia bezpieczeństwa, które wymagają specyficznego podejścia.

Kluczowe pytania

- Jakie nowe zagrożenia bezpieczeństwa wprowadza wykorzystanie AI?
- Czy posiadamy odpowiednie zabezpieczenia przed atakami specyficznymi dla AI?
- Jak chronimy dane wykorzystywane do trenowania i zasilania systemów AI?

Możliwe działania

- Przeprowadzenie oceny ryzyk cyberbezpieczeństwa związanych z AI.
- Rozszerzenie strategii cyberbezpieczeństwa o elementy specyficzne dla AI.
- Regularne testy penetracyjne systemów AI.

4.2.5. Nieautoryzowane użycie AI przez pracowników

Pracownicy często korzystają z narzędzi AI niezależnie od oficjalnych polityk, co stwarza zarówno szanse, jak i ryzyka dla organizacji, zwłaszcza w kwestii wycieku danych.

Kluczowe pytania

- Czy mamy świadomość, że nasi pracownicy mogą już korzystać z narzędzi AI na własną rękę?
- Jakie są potencjalne ryzyka (np. bezpieczeństwo danych, compliance) i korzyści (np. innowacje) takiego zachowania?
- Czy mamy wdrożone polityki i wytyczne dotyczące korzystania z zewnętrznych narzędzi AI?

Możliwe działania

- Przeprowadzenie ankiety lub audytu wykorzystania narzędzi AI przez pracowników.
- Opracowanie wytycznych bezpiecznego korzystania z publicznych narzędzi AI.
- Rozważenie wdrożenia wewnętrznych alternatyw dla popularnych narzędzi AI.

4.2.6. Przeciwdziałanie stronniczości algorytmów

Systemy AI mogą utrzymywać lub wzmacniać istniejące uprzedzenia, co wymaga świadomego przeciwdziałania.

Kluczowe pytania

- Jak identyfikujemy potencjalne stronniczości w naszych systemach AI?
- Jakie środki stosujemy, aby zapewnić neutralność i sprawiedliwość systemów AI?
- Czy testujemy nasze algorytmy pod kątem dyskryminacji różnych grup?

Możliwe działania

- Wdrożenie procedur testowania algorytmów pod kątem stronniczości.
- Ustanowienie różnorodnych zespołów odpowiedzialnych za rozwój i ocenę systemów AI.
- Regularne audyty wyników systemów AI pod kątem równego traktowania.



Przykład pytania

Jakie działania podejmujemy, aby unikać stronniczości w systemach AI?



Przykład działania

Wdrożenie standardów etycznych oraz systematycznych testów sprawdzających algorytmy pod kątem równości i niedyskryminacji.

4.2.7. Ochrona danych i prywatności

Wykorzystanie AI wiąże się z przetwarzaniem dużych ilości danych, co wymaga szczególnej uwagi w zakresie ochrony prywatności.

Kluczowe pytania

- Czy nasze wykorzystanie danych w systemach AI jest zgodne z regulacjami?
- Jak zapewniamy przejrzystość wykorzystania danych osobowych w AI?
- Czy stosujemy zasadę minimalizacji danych w kontekście AI?

Możliwe działania

- Przeprowadzenie oceny wpływu na ochronę danych dla inicjatyw AI.
- Wdrożenie technicznych mechanizmów ochrony prywatności w systemach AI.
- Ustanowienie procedur zarządzania zgodą na wykorzystanie danych w AI.

4.2.8. Polityki wewnętrzne

Jasne polityki regulujące wykorzystanie AI stanowią fundament odpowiedzialnego podejścia.

Kluczowe pytania

- Czy posiadamy kompleksowe polityki regulujące wykorzystanie AI?
- Jak zapewniamy, że pracownicy znają i rozumieją te polityki?
- Jak często aktualizujemy polityki w odpowiedzi na zmiany technologiczne?

Możliwe działania

- Opracowanie polityki wykorzystania AI w organizacji.
- Wdrożenie programu szkoleń z polityk AI dla pracowników.
- Ustanowienie procesu regularnej rewizji polityk AI.

4.2.9. Wyjaśnialność i transparentność

W wielu branżach lub konkretnych sytuacjach możliwość wyjaśnienia decyzji podejmowanych przez systemy AI jest kluczowa dla budowania zaufania i zgodności z regulacjami. Wyjaśnialność wyklucza pewną klasę algorytmów (np. sieci neuronowe) w niektórych zastosowaniach. Na przykład generowanie grafiki może nie być wyjaśnialne bo koszt „pomyłki” algorytmu jest akceptowalny. Niemniej, w innych zastosowaniach typu diagnoza medyczna, decyzja kredytowa, rekrutacja albo pojazd autonomiczny pojawia się poważna odpowiedzialność, konieczność wyjaśnienia jak doszło do podjęcia takiej a nie innej decyzji, a także wymóg monitorowania działania rozwiązania AI. W praktyce w takich zastosowaniach dostawcy nie powinni proponować algorytmów niewyjaśnialnych, nawet gdyby czasem dawały lepsze efekty. Stosowanie wyjaśnialnych modeli ma też dodatkową zaletę – umożliwia lepsze skonkretyzowanie działań na podstawie rekomendacji modelu. Przykładowo w

wypadku modeli segmentacji klientów spółka może bardziej precyzyjnie budować komunikację do danego segmentu klientów, jeśli jest w stanie go opisać i zrozumieć.

Z drugiej strony bardziej zaawansowane lub precyzyjne modele (np. oparte o sieci neuronowe) mogą dawać wyniki o wyższej jakości, choć niezrozumiałe dla człowieka. Odpowiedzialnością spółki jest zdecydować, w jakich sytuacjach priorytetyzować wyjaśnialność modeli, nawet potencjalnie kosztem jakości.

Kluczowe pytania

- W jakim stopniu jesteśmy w stanie zrozumieć, jak działają nasze systemy AI i jakie są podstawy ich decyzji?
- Czy wymóg wyjaśnialności jest uwzględniany przy wyborze i projektowaniu systemów AI?
- Czy możemy wytłumaczyć interesariuszom, jak działają nasze systemy AI?

Możliwe działania

- Ustanowienie standardów wyjaśnialności dla różnych zastosowań AI w organizacji.
- Wdrożenie narzędzi do śledzenia i dokumentowania decyzji systemów AI.
- Tam, gdzie to możliwe i wymagane, wybieranie systemów AI, których działanie jest zrozumiałe i wyjaśnialne.

4.2.10. Podsumowanie

Nadzór nad AI w firmie to proces wieloetapowy, łączący strategiczne planowanie, zarządzanie ryzykiem i rozwój kompetencji. Kluczowymi elementami są:

- Rzetelna ocena obecnej pozycji i wyznaczenie celów strategicznych
- Jasne przypisanie odpowiedzialności na szczeblu zarządu i rady nadzorczej
- Integracja AI z systemami kontroli, bezpieczeństwa i zarządzania talentem
- Stały monitoring zgodności rozwiązań z wewnętrznymi politykami oraz standardami etyki i ochrony danych

Rekomenduje się, aby w strukturach organizacyjnych powołać komitet ds. AI, który będzie regularnie przedstawiał radzie nadzorczej aktualizacje oraz wyniki badań wpływu wdrożonych rozwiązań na działalność firmy. Takie podejście umożliwi elastyczne reagowanie na zmieniające się warunki rynkowe oraz ewolucję technologii AI.

Systematyczne odnoszenie się do przedstawionych pytań i wdrażanie rekomendowanych działań pozwoli na świadome i odpowiedzialne wykorzystanie sztucznej inteligencji, minimalizując ryzyka i maksymalizując korzyści biznesowe.

Rekomendujemy członkom i członkiniom rad nadzorczych aktywne zaangażowanie w nadzór nad AI poprzez regularne zadawanie przedstawionych pytań zarządowi oraz monitorowanie postępów w realizacji rekomendowanych działań.

4.3. Ryzyka i Regulacje – monitorowanie otoczenia regulacyjnego

Rady nadzorcze muszą objąć nadzorem AI, która jest nowym obszarem o jednocześnie ogromnym potencjale transformacyjnym i ryzyku. Kluczowym elementem tego nadzoru jest śledzenie dynamicznie zmieniającego się otoczenia regulacyjnego. Poniższa sekcja zawiera przegląd podstawowych zagadnień prawnych związanych z danymi wykorzystywanymi przez AI, przegląd obowiązujących regulacji w UE oraz sytuację w Polsce.

4.3.1. Przegląd obecnych regulacji dotyczących AI

Sztuczna inteligencja w coraz większym stopniu wpływa na modele biznesowe, procesy operacyjne i decyzje zarządcze w firmach. W odpowiedzi na te zmiany Unia Europejska konsekwentnie rozwija ramy regulacyjne, których celem jest zapewnienie bezpiecznego, etycznego i odpowiedzialnego wykorzystania AI.

AI ukierunkowana na człowieka – filar unijnego podejścia

Unia Europejska promuje model AI ukierunkowanej na człowieka (human-centric AI), który oznacza, że systemy sztucznej inteligencji powinny:

- wspierać wolność jednostki i prawa człowieka,
- być przejrzyste, możliwe do audytu i kontrolowane przez człowieka,
- nie podejmować decyzji autonomicznie bez odpowiedzialności ludzkiej,
- sprzyjać sprawiedliwości i zapobiegać dyskryminacji.

Kontekst regulacyjny w UE – podejście systemowe i prewencyjne

Unia Europejska przyjęła unikalne podejście do regulacji AI, oparte na:

- klasyfikacji ryzyka (np. AI Act),
- ochronie praw podstawowych (np. poprzez RODO),
- zapewnieniu przejrzystości i zaufania (np. poprzez DSA, DMA),
- otwartości na wymianę danych (np. Data Act, DGA).

Kluczowe regulacje UE dotyczące AI

- AI Act (2024) – klasyfikacja systemów wg ryzyka, obowiązki zgodności, audyty, sankcje do 7% globalnego obrotu.
- RODO (GDPR) – zasady przetwarzania danych osobowych, z nowymi wytycznymi nt. AI i profilowania.
- NIS2 – cyberbezpieczeństwo systemów AI, obowiązek zgłaszania incydentów.
- DSA (Digital Services Act) – obowiązki dot. przejrzystości algorytmów i moderacji treści dla VLOPs (Very Large Online Platforms – >45 mln użytkowników w UE).
- DMA (Digital Markets Act) – regulacje dla „gatekeepers” (Google, Apple, Amazon etc.) – transparentność, interoperacyjność.

- Data Act (2023) – dostępność i ponowne wykorzystanie danych z urządzeń IoT.
- Data Governance Act (DGA) – bezpieczne mechanizmy udostępniania danych publicznych i przemysłowych.
- Dyrektywa DSM (2019/790/UE) – wyjątki dla text and data mining, istotne dla trenowania modeli AI.
- Dyrektywa 96/9/WE – ochrona baz danych (prawo *sui generis*).
- Standardy ISO:
 - ISO/IEC 42001 – system zarządzania AI,
 - ISO 23894 – zarządzanie ryzykiem w AI.

TABELA Porównanie regulacji AI

Regulacja	Główne założenia	Wpływ na spółki	Typowe spółki	Rekomendacje dla RN
AI Act	Klasyfikacja wg ryzyka, obowiązki zgodności	Nowe procesy, dokumentacja, audyty	Firmy używające AI decyzyjnie (finanse, HR, zdrowie)	Identyfikacja systemów AI, ocena ryzyka, plan zgodności
RODO	Ograniczenia dla automatycznego podejmowania decyzji, zgoda, informowanie	Kary za brak zgody przy trenowaniu modeli	Wszystkie firmy przetwarzające dane osobowe	Dokumentacja zgód, anonimizacja, polityki danych
NIS2	Zgłaszanie incydentów, audyty cyberbezpieczeństwa	Wymogi bezpieczeństwa systemów AI	Finanse, energetyka, infrastruktura	Audyty bezpieczeństwa AI, osoby odpowiedzialne
DSA	Przejrzystość algorytmów, raportowanie	Dotyczy tylko VLOPs (>45 mln użytkowników w UE)	Meta, Google, TikTok, Amazon	Dokumentuj i ujawniaj użycie AI w moderacji
DMA	Zakaz nadużyć przez największe platformy	Obowiązki interoperacyjności, zakaz faworyzowania	Gatekeepers (Google, Apple itd.)	Śledź obowiązki dostawców AI, dostosuj integracje
Data Act	Otwartość danych z IoT	Nowe obowiązki udostępniania danych	Producenci IoT, przemysł, auta	Określenie danych podlegających udostępnieniu
DGA	Ponowne wykorzystanie danych sektora publicznego	Konieczność zgodności i dokumentacji	Administracja, przemysł danych	Weryfikacja zobowiązań dot. danych publicznych
Dyrektywa DSM (2019/790/UE)	TDM bez zgody dla nauki, z opt-out dla komercji	Legalność scrapowania zależy od zabezpieczeń	AI, wydawnictwa, media	Weryfikuj legalność pozyskiwanych danych
Dyrektywa 96/9/WE	Ochrona struktur baz danych przy znacznych nakładach	Możliwość ochrony własnych zbiorów danych	E-commerce, instytucje, big data	Dokumentuj strukturę i nakłady baz danych
ISO/IEC 42001	Zarządzanie AI – zgodność, transparentność	Standaryzacja podejścia do AI governance	Wszystkie firmy wdrażające AI	Wdrożenie jako baza do zgodności z AI Act

Regulacja	Główne założenia	Wpływ na spółki	Typowe spółki	Rekomendacje dla RN
ISO 23894	Ryzyko w AI – identyfikacja i ocena	Możliwość systematycznego zarządzania ryzykiem	Finanse, medycyna, ubezpieczenia	Stwórz mapę ryzyk, integruj z nadzorem wewnętrznym

Case studies i orzecznictwo

LG Hamburg 310 O 221/23 (2024)

Sąd orzekł, że wykorzystanie chronionych prawem autorskim treści z witryn internetowych (należących m.in. do Heise Medien) przez OpenAI bez zgody twórcy do trenowania modeli AI narusza prawo autorskie. To jedno z pierwszych orzeczeń tego typu w Europie, pokazujące, że „scraping” treści do trenowania AI może być nielegalny nawet w przypadku publicznie dostępnych stron.

Pokazuje to, że nawet treści publicznie dostępne są chronione, jeśli właściciel zabezpieczył prawa. Ten wyrok ma implikacje zarówno dla portali publikujących treści (wystarczy odpowiednie zabezpieczenia prawne, by chronić treści), jak i dla stron chcących korzystać z publicznie dostępnych treści do trenowania modeli (należy zweryfikować ich zabezpieczenia prawne)

Źródło: <https://www.lto.de/recht/hintergruende/h/lg-hamburg-openai-heise-medien-urheberrecht-kuenstliche-intelligenz-ki-scraping-training/>

Clearview AI (UE, 2021–2022)

Firmie zakazano dalszego przetwarzania danych biometrycznych mieszkańców UE bez ich zgody. Wyrok pokazuje, że nawet firmy z USA muszą respektować unijne zasady prywatności i transparentności, gdy ich AI operuje na danych obywateli UE.

Firmy spoza UE muszą przestrzegać europejskich zasad ochrony danych.

Replika.ai – Włochy (2023)

Włoski organ ochrony danych zakazał działalności chatbotowi Replika.ai na terenie Włoch z uwagi na brak wystarczających zabezpieczeń chroniących dane osobowe i interakcje z użytkownikami – co miało znaczenie przy ocenie legalności AI pod kątem RODO.

AI musi mieć mechanizmy zgodne z RODO i zabezpieczenia dla grup wrażliwych.

Uber (USA, 2018)

Wypadek śmiertelny z udziałem autonomicznego pojazdu Ubera doprowadził do śledztwa i podważenia procesu testowania i nadzoru nad systemem AI. Z punktu widzenia regulacji AI – przykład znaczenia odpowiedzialności produktowej.

Znaczenie testów, audytów bezpieczeństwa i nadzoru człowieka nad AI.

OpenAI vs FTC (USA, 2023)

Federalna Komisja Handlu w USA wszczęła postępowanie przeciw OpenAI w związku z możliwością generowania nieprawdziwych i szkodliwych informacji przez modele językowe oraz brakiem transparentności w zakresie źródeł danych.

Organy regulacyjne zaczynają stosować ochronę konsumentką do AI.

Lista kontrolna zgodności z prawem – pytania do zadania przez rady nadzorcze

- Czy systemy AI są sklasyfikowane pod kątem AI Act?
- Czy dane treningowe są zgodne z RODO (zgody, anonimizacja)?
- Czy wdrożono framework zarządzania ryzykiem AI (ISO/NIST)?
- Czy znane są obowiązki pod Data Act lub DGA?
- Czy osoby odpowiedzialne za AI są wyznaczone i przeszkolone?

Benchmark międzynarodowy – UE, USA, Chiny

Porównanie podejść do regulacji AI w różnych jurysdykcjach:

Aspekt	UE	USA	Chiny
Podejście	Prewencyjne, prawa człowieka	Wolnorynkowe, sektorowe	Centralne, kontrolne
Główna regulacja	AI Act + GDPR	Brak – wytyczne NIST, FTC	Obowiązkowe zatwierdzanie algorytmów
Priorytety	Przejrzystość, odpowiedzialność	Innowacja, elastyczność	Nadzór, moderacja treści

Praktyczne zalecenia zarządcze dla spółek i rad nadzorczych

- Utworzenie mapy systemów AI stosowanych w firmie i ich klasyfikacji ryzyka (AI Act).
- Przeprowadzenie audytu danych wejściowych do trenowania modeli – ich legalności, zgodności z RODO.
- Wyznaczenie zespołu AI governance – osoby odpowiedzialne za zgodność, etykę, nadzór.
- Wdrożenie framework'u zarządzania ryzykiem AI, np. ISO 23894 lub NIST AI RMF.
- Integracja zgodności z regulacjami AI z systemem kontroli wewnętrznej i audytami IT.
- Wprowadzenie scenariuszy reagowania na naruszenia związane z AI, np. błędne decyzje algorytmu, wycieki danych.

Źródła wiedzy o zmianach regulacyjnych

- Oficjalne strony Komisji Europejskiej: <https://digital-strategy.ec.europa.eu>
- OECD AI Policy Observatory: <https://oecd.ai>

- EU AI Act Tracker: <https://artificialintelligenceact.eu>
- Future of Life Institute: <https://futureoflife.org>
- Standards ISO: <https://www.iso.org/committee/6794475.html>
- European Data Protection Board (EDPB): <https://edpb.europa.eu>

Podsumowanie – główne wnioski dla rad nadzorczych

- Zgodność z przepisami AI to nowy wymiar odpowiedzialności zarządczej i nadzorczej – wpływa na reputację, bezpieczeństwo, strategię danych.
- Rady nadzorcze muszą aktywnie uczestniczyć w nadzorze nad wdrażaniem AI, nie tylko jako kwestii technologicznej, ale także strategicznej i zgodności.
- Nowe przepisy mają zastosowanie także do mniejszych firm, które wdrażają AI w krytycznych procesach (HR, finanse, zdrowie).
- Przewaga rynkowa może wynikać z etycznego, transparentnego i zgodnego wdrożenia AI – szczególnie wobec partnerów i inwestorów.
- Ważne do wykonania kroki w procesie nadzoru:
 - Zidentyfikowanie ryzyk regulacyjnych i cyberbezpieczeństwa związanych z AI.
 - Monitorowanie regulacji i aktywne wdrażanie nowych obowiązków.
 - Zapewnienie edukacji i szkoleń dla zarządów i pracowników dotyczących zgodności z regulacjami.
 - Regularne audyty zgodności i zarządzanie ryzykiem związanym z AI.

4.3.2. Co jest chronione prawem, a co nie?

Kontekst: dane jako nowy zasób strategiczny

Dane stały się jednym z kluczowych czynników wzmacniania konkurencyjności współczesnych przedsiębiorstw. W dobie transformacji cyfrowej i wdrażania narzędzi opartych o sztuczną inteligencję, jakość danych, ich dostępność oraz legalność ich wykorzystania decydują nie tylko o jakości modeli AI, ale także o przewadze strategicznej. Modele trenowane na niskiej jakości danych generują wyniki obciążone błędami, podczas gdy dobrze zarządzana infrastruktura danych stanowi podstawę wiarygodnych i bezpiecznych rozwiązań AI.

W tym kontekście aspekty prawne związane z wykorzystaniem danych nabierają nowego znaczenia. Legalność pozyskiwania danych, ich przetwarzania oraz ochrona własnych zbiorów staje się strategicznym tematem dla rad nadzorczych. Bardzo istotna jest wiedza o tym, co jest chronione prawem i w jaki sposób – bazy danych, zawartość stron internetowych, dane wykorzystywane do trenowania modeli.

Ochrona prawna surowych danych, baz danych i treści twórczych

Czy same dane są chronione prawem autorskim?

Nie.

Same dane, fakty (np. „temperatura w Warszawie wynosi 27°C” lub zapis poziomu temperatury w kolejnych dniach) nie podlegają ochronie prawa autorskiego. Są to informacje obiektywne, a nie rezultat twórczości. Oznacza to, że jeśli strona trzecia wejdzie w posiadanie takich danych zebranych przez spółkę, może z nich korzystać, trenować na ich podstawie modele AI bądź wykorzystywać w inny sposób.

„Dane jako takie nie podlegają ochronie prawa autorskiego” – wyrok TSUE w sprawie C-545/07.

Jednak:

- układ danych w uporządkowaną strukturę (np. bazę danych) może podlegać ochronie na mocy tzw. prawa **sui generis** (**Dyrektywa 96/9/WE**);
- treści inne niż dane zawarte na stronach internetowych mogą być chronione prawem autorskim (np. opisy produktów, grafiki, artykuły).

Dopiero **strukturalne opracowanie danych** (np. baza danych) lub **twórcze przedstawienie** (np. opracowania, komentarze, wizualizacje) mogą być chronione na gruncie prawa.

Kiedy chroniona jest baza danych?

Baza danych może być chroniona na dwa sposoby:

a) Prawem autorskim (jeśli ma twórczy układ lub dobór danych):

- Dotyczy sytuacji, gdy dane zostały twórczo opracowane, pogrupowane, zestawione w oryginalny sposób.
- Ochrona powstaje jak dla każdego utworu: automatyczna, bez rejestracji.

b) Prawem *sui generis* (ochrona inwestycji w bazę danych):

- Dotyczy baz nietwórczych, jeśli ich opracowanie wymagało istotnych nakładów (czasowych, finansowych, organizacyjnych).
- Zapewnia ochronę przed kopiowaniem całości lub istotnych części.
- Podstawa prawna: Dyrektywa 96/9/WE o ochronie baz danych.

Metody ochrony danych

Czy baza danych jest chroniona „automatycznie”?

To zależy od rodzaju ochrony:

OCHRONA PRAWEM AUTORSKIM

- Działa automatycznie, ale tylko jeśli baza danych ma twórczy charakter, czyli sposób doboru, zestawienia, układu danych odzwierciedla indywidualne decyzje twórcze (np. autorska klasyfikacja, struktura).

W przypadku braku tych cech, ochrona autorska nie przysługuje.

Dotyczy również nietwórczych baz danych, jednak wymagane jest wykazanie, że ich stworzenie wiązało się ze znacznym nakładem:

- finansowym (np. zakup danych),
- czasowym (analiza, strukturyzacja),
- organizacyjnym (budowa systemu, utrzymanie aktualności).

Brak udokumentowania lub sygnalizacji ochrony może skutkować brakiem możliwości skutecznego egzekwowania praw. W takiej sytuacji strony trzecie mogą wykorzystać informacje zamieszczone na stronie internetowej do trenowania własnych modeli i ograniczyć przewagę konkurencyjną spółki.

Ochrona treści na stronach internetowych

Czy istotne jest zaznaczenie ochrony na stronie internetowej?

Zdecydowanie tak.

Nawet jeśli ochrona przysługuje na mocy prawa, brak technicznych i prawnych sygnałów może uniemożliwić skuteczne egzekwowanie roszczeń.

Dlatego spółka powinna:

- dokumentować wkład w opracowanie bazy,
- zamieścić komunikat: „wszelkie prawa zastrzeżone”,
- wskazać, że baza lub układ danych podlega ochronie i z jakiego tytułu,
- zastosować plik `robots.txt` z dyrektywami `noai`, `noindex` oraz nagłówki HTTP (`X-Robots-Tag`), które wskazują jednoznacznie, że właściciel strony nie zgadza się na użycie danych do trenowania modeli,
- przygotować regulamin (tzw. **Terms & Conditions** – T&C).

Czy treści na stronie internetowej są chronione?

Tak – jeśli spełniają warunki twórczości.

- Teksty, zdjęcia, artykuły podlegają prawu autorskiemu, jeśli są oryginalne.
- Jednak w przypadku braku zabezpieczeń i cech twórczych, treści mogą zostać legalnie wykorzystane, np. do trenowania modeli AI.

Dostępność danych do trenowania modeli AI (Text & Data Mining)

Text and Data Mining (TDM) oznacza zautomatyzowaną analizę dużych zbiorów danych i tekstów, wykorzystywaną m.in. do trenowania modeli AI. Może obejmować analizę treści tekstowych, zdjęć, dźwięków czy danych pomiarowych.

Kontekst prawny – Unia Europejska:

Dyrektywa DSM (2019/790/UE) wprowadza dwa wyjątki umożliwiające legalne prowadzenie TDM:

DLA CELÓW BADAWCZYCH (NON-COMMERCIAL):

- Dozwolone jest trenowanie modeli AI na danych upublicznionych przez instytucje badawcze oraz biblioteki cyfrowe – bez konieczności uzyskania zgody właściciela treści. (Art. 3 dyrektywy DSM).

DLA CELÓW KOMERCYJNYCH:

- Możliwe jest wykorzystywanie danych tylko wtedy, gdy właściciel treści nie zastrzegł tzw. opt-outu (Art. 4 dyrektywy DSM).
- Opt-out musi być wyraźnie zasygnalizowany, np. poprzez:
- zamieszczenie pliku ``robots.txt`` z dyrektywą ``noai``,
- nagłówki HTTP (``X-Robots-Tag: noai``),
- jasne oznaczenie na stronie internetowej,
- zapis w regulaminie (Terms & Conditions).

Stanowisko Polski (2024):

- Polska nadal nie wdrożyła dyrektywy DSM do krajowego porządku prawnego – choć termin minął w czerwcu 2021 roku.
- W 2024 roku opublikowano projekt nowelizacji ustawy o prawie autorskim, który dopiero teraz implementuje TDM, wzorując się na rozwiązaniach unijnych.
- Projekt zakłada wprowadzenie obu wyjątków (Art. 3 i Art. 4 DSM) – jednak niektóre środowiska postulują szerszy zakaz wykorzystywania treści w AI.
- Do czasu uchwalenia przepisów krajowych, obowiązuje bezpośrednie stosowanie przepisów unijnych w interpretacji sądów i praktyce.
- Źródło: Z. Okoń, TDM, [sztuczna inteligencja i projekt nowelizacji prawa autorskiego](#) (LinkedIn):

Dlaczego to ważne z perspektywy rad nadzorczych?

- Spółka może posiadać istotne zasoby danych i treści, które – bez odpowiedniego zabezpieczenia – mogą zostać legalnie wykorzystane przez podmioty trzecie.
- Brak wdrożonego opt-outu może oznaczać utratę kontroli nad wartością informacyjną i przewagą konkurencyjną.
- Jednocześnie nowe przepisy tworzą możliwości trenowania modeli AI na cudzych, publicznych danych, które nie zostały zabezpieczone – co może stanowić szansę strategiczną.
- Rady nadzorcze powinny nadzorować zarówno proces zabezpieczania własnych danych, jak i analizę potencjalnych źródeł danych dostępnych publicznie.

Sposoby zabezpieczania danych i treści

Jakie zabezpieczenia należy wdrożyć na stronach internetowych?

Techniczne:

- ``robots.txt`` z dyrektywami ``noai``, ``noindex``,
- nagłówki HTTP: ``X-Robots-Tag: noai, noindex``,
- blokowanie znanych adresów IP firm AI.

Prawne i regulacyjne:

- komunikat: **, „wszelkie prawa zastrzeżone” **,
- regulamin (T&C),
- dokumentacja inwestycji w bazę danych,
- opcjonalna rejestracja bazy danych.

Kto ponosi odpowiedzialność – właściciel danych czy podmiot wykorzystujący?

Obie strony mają określone obowiązki:

WŁAŚCICIEL DANYCH POWINIEN:

- zabezpieczyć dane (opt-out, regulamin, metadane),
- sygnalizować przysługujące prawa.

PODMIOT WYKORZYSTUJĄCY DANE ZOBOWIĄZANY JEST DO:

- sprawdzenia czy dane są zastrzeżone,
- respektowania sygnałów opt-out,
- niewykorzystywania chronionych treści bez zgody.

Jeśli dane są upublicznione, niezabezpieczone prawnie i nie mają cech twórczych, mogą być legalnie wykorzystywane do TDM.

Dane z urządzeń IoT – zmiany w Data Act

Zgodnie z Data Act (2023/2854), który wejdzie w życie w 2025 roku, przedsiębiorstwa produkujące urządzenia generujące dane (np. samochody, urządzenia IoT, smartwatche) będą zobowiązane udostępniać zebrane dane użytkownikom oraz (na żądanie) stronom trzecim w ustandaryzowany i otwarty sposób. Ma to na celu uniknięcie monopolizacji dostępu do danych i ułatwienie rozwoju nowych usług oraz AI opartych na rzeczywistych danych.

Źródło: [Rozporządzenie \(UE\) 2023/2854 w sprawie danych:](#)

Obowiązki w ramach Data Act:

- umożliwienie użytkownikom dostępu do danych generowanych przez urządzenia podłączone do internetu (IoT),
- udostępnienie tych danych stronom trzecim na żądanie.

Dlaczego to ważne?

- otwiera to rynek dla podmiotów chcących trenować AI na rzeczywistych danych (np. dane z pojazdów, smart zegarków),
- dla firm będących producentami sprzętu to ryzyko utraty monopolu informacyjnego,
- dla innych – szansa na pozyskanie danych bez kosztów licencyjnych.

Jak się przygotować?

- zidentyfikować dane generowane przez własne produkty,
- wdrożyć kontrolę nad strukturą udostępniania danych,

- przygotować politykę udzielania dostępu i model licencyjny.

Dane osobowe i trenowanie AI – ryzyka RODO

Dane osobowe należą do najcenniejszych zasobów danych w organizacji. Odpowiednio przetwarzane i analizowane mogą wspierać rozwój modeli AI, pozwalając m.in. na personalizację usług, tworzenie lepszych interfejsów z klientem, dopasowanie ofert czy przewidywanie potrzeb.

Jednak ich wykorzystanie wymaga zachowania zgodności z Rozporządzeniem o Ochronie Danych Osobowych (GDPR/RODO) – szczególnie w kontekście trenowania AI na danych użytkowników.

Trenowanie modeli AI na danych osobowych wymaga spełnienia warunków RODO:

- określenie podstawy prawnej (np. zgoda, umowa, uzasadniony interes),
- obowiązek informacyjny wobec osób, których dane dotyczą,
- zapewnienie możliwości sprzeciwu i kontroli nad przetwarzaniem danych.

Przykłady naruszeń:

- trenowanie chatbotów na archiwalnych rozmowach klientów bez ich zgody,
- analiza nagrań głosowych bez wskazania celu i odbiorców danych.

Rekomendacje dla rad nadzorczych:

- wymagać audytu źródeł danych wykorzystywanych do trenowania modeli,
- weryfikować istnienie dokumentacji zgód lub uzasadnień prawnych,
- kontrolować, czy dane są anonimizowane lub pseudonimizowane.

Dlaczego to ważne z perspektywy rady nadzorczej?

- Dane osobowe to zasób strategiczny, który – przetwarzany zgodnie z prawem – może zwiększać przewagę konkurencyjną firmy.
- Modele AI oparte na rzeczywistych danych klientów umożliwiają lepsze zrozumienie użytkownika, automatyzację procesów i personalizację produktów/usług.
- Rada nadzorcza powinna jednak równoważyć korzyści biznesowe z wykorzystania danych z potencjalnym ryzykiem prawnym i reputacyjnym.

Przykłady zgodnego z prawem wykorzystania danych osobowych w AI:

- systemy rekomendacyjne oparte na historii zakupów klienta (na podstawie zgody),
- automatyzacja odpowiedzi na pytania klientów na podstawie zanonimizowanych interakcji z infolinią,
- segmentacja użytkowników pod kątem potrzeb i preferencji – przy odpowiednim obowiązku informacyjnym i podstawie przetwarzania.

Podsumowanie i rekomendacje dla rad nadzorczych

Podsumowanie kluczowych wniosków

- Pojedyncze dane nie podlegają ochronie autorskiej, ale ich strukturalne opracowanie (np. baza danych) lub twórcze przedstawienie (np. komentarze, wizualizacje) może podlegać ochronie.
- Baza danych może być chroniona:
 - prawem autorskim (jeśli ma twórczy układ),
 - prawem *sui generis* (jeśli wiąże się z istotnym nakładem).
- Treści na stronach internetowych są chronione, jeśli spełniają przesłanki twórczości. Niezabezpieczone, nietwórcze treści mogą być legalnie wykorzystywane (np. do trenowania AI).
- Przepisy pozwalają trenować modele AI na danych z internetu, jeśli właściciel nie zgłosił sprzeciwu (opcja opt-out).
- Bez wdrożonych zabezpieczeń technicznych i prawnych, egzekwowanie ochrony praw może być nieskuteczne.
- Od 2025 r. Data Act zobowiązuje producentów urządzeń IoT do udostępniania danych użytkownikom i stronom trzecim, co może wpłynąć na przewagi informacyjne firm.
- Trenowanie AI na danych osobowych wymaga zgodności z RODO – m.in. uzyskania zgody, spełnienia obowiązku informacyjnego oraz umożliwienia sprzeciwu.

Rekomendacje dla rad nadzorczych

- Zapewnić, że spółka posiada strategię zarządzania danymi i treściami – uwzględniającą ochronę, wykorzystywanie i udostępnianie.
- Wymagać przeglądu zawartości stron i systemów, które mogą być źródłem danych do trenowania AI (zarówno używanych, jak i udostępnianych przez spółkę).
- Zatwierdzić politykę zabezpieczeń technicznych i regulaminowych (np. T&C, robots.txt).
- Upewnić się, że spółka dokumentuje inwestycje w tworzenie baz danych – w celu ewentualnej ochrony *sui generis*.

Pytania kontrolne do zarządu

W zakresie ochrony danych i treści:

- Czy wiadome jest, które dane, treści i bazy są publicznie udostępniane?
- Czy strona internetowa zawiera zabezpieczenia przed scrapowaniem i trenowaniem AI (np. robots.txt, nagłówki, T&C)?
- Czy wiadomo, które dane zostały wykorzystane do trenowania modeli AI? Czy ich wykorzystanie jest legalne i zgodne z regulacjami?
- Czy spółka posiada dokumentację wskazującą na wkład inwestycyjny w bazy danych?

W kontekście trenowania AI:

- Czy dane wykorzystywane do trenowania modeli są pozyskane zgodnie z prawem (TDM, opt-out, zgody)?
- Czy stosowane są mechanizmy opt-out chroniące treści przed wykorzystaniem przez podmioty trzecie (np. firmy AI)?

W zakresie danych IoT i zgodności z Data Act:

- Czy spółka zidentyfikowała dane generowane przez urządzenia, które podlegają obowiązkowi udostępniania?

- Czy istnieje polityka dotycząca udzielania dostępu do tych danych oraz model ich licencjonowania?

W zakresie danych osobowych i zgodności z RODO:

- Czy dane osobowe wykorzystywane do trenowania modeli są zgodne z wymaganiami RODO (m.in. zgody, podstawy prawne, informowanie)?
- Czy spółka zapewnia możliwość sprzeciwu wobec przetwarzania oraz czy dane są odpowiednio pseudonimizowane lub anonimizowane?

4.3.3. Sytuacja w Polsce – wdrażanie regulacji i programy wsparcia

Kontekst dla rad nadzorczych:

W Polsce intensywnie wdrażane są regulacje unijne dotyczące sztucznej inteligencji. Rozumienie lokalnego kontekstu regulacyjnego oraz wykorzystanie dostępnych programów wsparcia umożliwia efektywniejsze zarządzanie zmianami i zapewnienie zgodności z regulacjami.

Stan implementacji regulacji UE w Polsce:

- AI Act wymaga ustanowienia krajowego organu nadzorczego – obecnie trwa proces legislacyjny,
- strategia AI wdrażana przez Ministerstwo Cyfryzacji obejmuje wykorzystanie AI w sektorze publicznym i prywatnym.

Potencjalne organy nadzoru AI w Polsce:

- Urząd Ochrony Danych Osobowych (UODO),
- Ministerstwo Cyfryzacji,
- Urząd Komunikacji Elektronicznej (UKE).

Programy wsparcia dla przedsiębiorstw (2024–2025):

- fundusze europejskie (FENG, Horyzont Europa) – finansowanie projektów AI zgodnych z regulacjami unijnymi,
- program GovTech Polska – promowanie innowacji technologicznych we współpracy z sektorem prywatnym,
- Rada ds. Sztucznej Inteligencji przy Ministerstwie Cyfryzacji – wsparcie merytoryczne i konsultacyjne dla przedsiębiorstw.

Rekomendacje dla rad nadzorczych w Polsce:

- aktywne uczestnictwo w programach wsparcia,
- monitorowanie postępów w implementacji regulacji,
- powołanie dedykowanych stanowisk compliance AI,
- regularne aktualizacje strategii AI oraz jej zgodności z regulacjami.

5. Co dalej? Zintegrowanie tematu AI w bieżącym procesie nadzoru

W miarę jak wykorzystanie systemów sztucznej inteligencji staje się coraz powszechniejsze, rady nadzorcze stoją przed wyzwaniem trwałego zintegrowania nadzoru nad AI z bieżącą praktyką ładu korporacyjnego. Ten rozdział przedstawia kluczowe obszary, na które powinny zwrócić uwagę rady, oraz sugeruje konkretne kroki, które mogą podjąć, aby skutecznie zarządzać tym procesem.

5.1. Normalizacja kwestii AI

Dlaczego to ważne

AI przestaje być domeną eksperymentów czy projektów pilotażowych – staje się integralną częścią strategii biznesowych i operacyjnych.

Działania dla rad nadzorczych:

- Umieszczanie tematyki AI jako stałego punktu w porządku obrad rady i komitetów.
- Wprowadzenie cyklicznych sesji informacyjnych z udziałem ekspertów wewnętrznych i zewnętrznych.
- Opracowanie ram komunikacyjnych dotyczących AI w kontekście strategii, ryzyk i odpowiedzialności.

Przykład:

Duże instytucje finansowe, takie jak JPMorgan, regularnie informują radę o postępach i ryzykach związanych z wykorzystaniem AI w zarządzaniu ryzykiem kredytowym i wykrywaniu oszustw.

5.2. Dostosowanie struktury rady nadzorczej

Dlaczego to ważne:

AI wpływa na strategiczne decyzje i wymaga nadzoru multidyscyplinarnego – często wykraczającego poza tradycyjne kompetencje.

Działania dla rad nadzorczych:

- Ustanowienie komitetu ds. technologii lub przegląd mandatu istniejącego komitetu (np. audytu, strategii, ryzyka).
- Rozważenie powołania administratora lub eksperta ds. AI jako doradcy rady.
- Przeprowadzenie oceny kompetencji technologicznych obecnych członków rady.

Przykład:

Niektóre przedsiębiorstwa, np. w firmach technologicznych notowanych na Nasdaq, uzupełniają skład rady lub zarządu o osoby z doświadczeniem w AI lub w transformacji cyfrowej.

5.3. Ocena kadry menedżerskiej w spółce

Dlaczego to ważne:

Zarząd i kadra kierownicza odgrywają kluczową rolę w operacjonalizacji strategii AI i zarządzaniu ryzykiem.

Działania dla rad nadzorczych:

- Uwzględnianie kompetencji cyfrowych i podejścia do AI w procesach oceny i nominacji kadry kierowniczej.
- Zadawanie pytań o integrację AI z procesami operacyjnymi, HR i zarządzaniem ryzykiem.
- Monitorowanie skuteczności wdrażania polityk dotyczących etycznego wykorzystania AI.

Przykład:

Firmy takie jak Unilever włączają kompetencje cyfrowe do corocznych ocen liderów.

5.4. Planowanie budżetu i planów na długi termin

Dlaczego to ważne:

AI wymaga inwestycji w infrastrukturę danych, kompetencje i bezpieczeństwo – decyzje budżetowe podejmowane dziś będą kształtować zdolność organizacji do adaptacji w przyszłości.

Działania dla rad nadzorczych:

- Analiza alokacji środków na rozwój AI, cyberbezpieczeństwo i zarządzanie danymi.
- Zapewnienie, że długoterminowe plany strategiczne uwzględniają transformację cyfrową z perspektywy AI.
- Ocenianie zasadności inwestycji nie tylko przez pryzmat ROI, ale także budowy kompetencji i odporności organizacji.

Przykład:

Zgodnie z badaniem EY, ponad 30% CEO planuje przesunąć środki z innych budżetów inwestycyjnych na rozwój AI.

5.5. Ofensywa zamiast defensywy – co dla każdego

Dlaczego to ważne:

Bierna postawa wobec AI może prowadzić do utraty konkurencyjności i zwiększenia ekspozycji na ryzyka regulacyjne lub reputacyjne.

Działania dla rad nadzorczych:

- Zdefiniowanie proaktywnej roli rady w inicjowaniu tematów AI – np. przegląd modeli biznesowych pod kątem transformacji z udziałem AI.
- Zachęcanie do eksperymentowania z AI w kontrolowanych warunkach.
- Kultura ciągłego doskonalenia i gotowości na zmiany.

Przykład:

Rada jednej z firm produkcyjnych w Europie Centralnej zainicjowała warsztaty strategiczne poświęcone wykorzystaniu AI w łańcuchu dostaw.

5.6. Ciągłe uczenie się – jak to może wyglądać

Dlaczego to ważne:

Technologia zmienia się szybciej niż tradycyjne cykle planowania i edukacji. Rady muszą zatem przestawić się na model permanentnego rozwoju kompetencji.

Działania dla rad nadzorczych:

- Organizowanie regularnych sesji szkoleniowych i „learning journeys” dla rady.
- Tworzenie „radarów trendów” i mechanizmów wczesnego ostrzegania we współpracy z działami innowacji lub partnerami zewnętrznymi.
- Stosowanie autoewaluacji wiedzy rady w obszarze AI oraz benchmarków branżowych.

Przykład:

W jednym z francuskich banków wprowadzono platformę e-learningową dla członków rady z modułami poświęconymi AI i ryzykom cyfrowym.

Zintegrowanie nadzoru nad AI z bieżącą pracą rady nadzorczej to proces wymagający konsekwencji, otwartości na zmiany i świadomego przywództwa. Dla rad, które podejmą to wyzwanie, AI może stać się nie tylko źródłem ryzyka, ale przede wszystkim motorem innowacji i przewagi strategicznej.

6. Poziomy integracji AI ze strategią firmy – przykłady możliwych obszarów w firmach oraz stopni integracji

Sztuczna inteligencja przestała być jedynie wizją przyszłości – dziś to kluczowy czynnik wzrostu i innowacji, pozwalający firmom osiągać wymierne korzyści. Nakłady, ryzyka i korzyści różnią się na każdym opisanym etapie wdrożenia AI. Celem tej części dokumentu jest opisanie możliwych zastosowań w podziale na poziomy zaawansowania wdrożenia AI wraz z przykładami przedsiębiorstw z Polski, Europy i całego świata.

6.1. Poziom 1 – wdrażanie AI w codziennych zadaniach

Cel: 10–20% wzrostu wydajności.

Na tym etapie AI automatyzuje rutynowe, powtarzalne zadania oraz wspiera pracowników inteligentnymi narzędziami.

Przykłady wdrażania AI w takich codziennych zadaniach w Europie:

Pierwszym przykładem może być branża Robotic Process Automation (RPA). Przykładową firmą z tej branży może być UiPath – globalny lider RPA automatyzujący procesy finansowo-księgowe (fakturowanie, wprowadzanie danych). Tego typu firma może wykorzystywać AI do:

- **Inteligentnej automatyzacji zadań (Hyperautomation):** łączenia tradycyjnej RPA z możliwościami AI, takimi jak uczenie maszynowe, przetwarzanie języka naturalnego (NLP), widzenie komputerowe i analiza predykcyjna, aby automatyzować bardziej złożone i oparte na danych procesy.
- **Ulepszania rozpoznawania i przetwarzania dokumentów (Intelligent Document Processing – IDP):** wykorzystania AI do inteligentnego wyodrębniania danych z różnorodnych dokumentów (np. faktur, umów), nawet tych nieustrukturyzowanych, co automatyzuje procesy związane z obiegiem dokumentów.
- **Wirtualnych asystentów opartych na AI:** umożliwiania tworzenia chatbotów i wirtualnych asystentów, którzy mogą wchodzić w interakcje z pracownikami i klientami, odpowiadać na pytania i wykonywać zadania.
- **Analizy procesów i identyfikacji możliwości automatyzacji (Process Mining i Task Mining):** wykorzystania AI do analizowania logów systemowych i interakcji użytkowników w celu identyfikowania wąskich gardeł, nieefektywności i potencjalnych obszarów do automatyzacji.
- **Inteligentnego rozplanowania narzędzi automatyzacji:** wykorzystania AI do dynamicznego zarządzania i optymalizacji pracy grup narzędzi automatyzacji (robotów softwarowych), zapewniając efektywne wykorzystanie zasobów i skalowalność automatyzacji.

- **Personalizacji doświadczeń użytkowników RPA:** dostarczania bardziej spersonalizowanych i intuicyjnych interfejsów oraz rekomendacji w procesie projektowania i wdrażania automatyzacji.
- **Predykcyjnej analizy i monitorowania wydajności robotów softwarowych:** wykorzystania AI do monitorowania pracy robotów, przewidywania potencjalnych problemów i optymalizowania ich wydajności.
- **Integracji z innymi platformami AI:** umożliwiania łatwej integracji z różnymi usługami i platformami AI innych dostawców, rozszerzając możliwości automatyzacji w przedsiębiorstwie.

<https://www.uipath.com/resources/automation-demo/job-matching-demo>

Kolejnym przykładem jest firma z branży projektowania i produkcji zaawansowanych rozwiązań inżynierskich w obszarach automatyki przemysłowej, energetyki, infrastruktury, transportu oraz cyfryzacji, np. Siemens. W tym przypadku AI może zostać wykorzystane do:

- **Zwiększania wydajności prac projektowych:** poprzez asystentów AI (Industrial Copilots) wspomagających projektowanie, generowanie kodu i zarządzanie dokumentacją techniczną.
- **Optymalizacji harmonogramów pracy:** wykorzystania AI do inteligentnego planowania i alokacji zasobów w zespołach inżynierskich i produkcyjnych.
- **Inteligentnej automatyzacji przemysłowej (Industrial AI):** łączenia AI z systemami automatyki w celu optymalizacji procesów produkcyjnych, kontroli jakości i konserwacji predykcyjnej.
- **Rozwoju Cyfrowych Bliźniaków (Digital Twins) wspomaganych przez AI:** tworzenia wirtualnych reprezentacji produktów i procesów, które dzięki AI umożliwiają symulacje, przewidywanie zachowań i optymalizację w świecie rzeczywistym.
- **Wspomagania w diagnostyce i konserwacji:** wykorzystania AI do analizy danych z maszyn i urządzeń w celu przewidywania awarii i optymalizacji strategii utrzymania.
- **Personalizacji i optymalizacji ofert dla klientów:** analizowania danych klientów w celu dostarczania bardziej dopasowanych rozwiązań i usług.
- **Wspierania zrównoważonego rozwoju:** wykorzystania AI do optymalizacji zużycia energii i zasobów w procesach przemysłowych.
- **Umożliwienia innowacji w projektowaniu produktów:** wykorzystania generatywnej AI do wspomagania inżynierów w tworzeniu nowych i ulepszonych projektów.
- **Poprawy jakości i skrócenia czasu weryfikacji oprogramowania:** wykorzystania AI w narzędziach EDA (Electronic Design Automation).

<https://itreseller.pl/siemens-prezentuje-przemyslowa-agentowa-sztuczna-inteligencje-zapowiedz-ery-przemyslu-5-0/>

Przykłady wdrażania AI w takich codziennych zadaniach w Polsce:

Platforma edukacyjna może wykorzystać AI do moderacji treści i spersonalizowanych rekomendacji, w celu usprawnienia obsługi użytkowników i zwiększenia efektywności zespołu. Firma taka jak **Brainly** może wykorzystać AI do:

- Zapewnienia wysokiej jakości i bezpieczeństwa treści.
- Dostarczania spersonalizowanych doświadczeń edukacyjnych.
- Ułatwienia dostępu do potrzebnych informacji.
- Zwiększenia efektywności operacyjnej.
- Tworzenia innowacyjnych narzędzi do nauki.
- Ciągłego ulepszania platformy w oparciu o dane.

Firma taka jak **DocPlanner (ZnanyLekarz)** – może wykorzystać AI w optymalizacji harmonogramów wizyt i automatycznych przypomnieniach dla pacjentów. Może to oznaczać wykorzystanie AI do:

- Ulepszenia jakości usług dla pacjentów poprzez personalizację i ułatwienie znalezienia odpowiedniej opieki.
- Optymalizacji pracy lekarzy i placówek medycznych poprzez inteligentne zarządzanie harmonogramami i wizytami.
- Zapewnienia wiarygodności i jakości informacji na platformie.
- Personalizacji komunikacji i działań marketingowych.
- Dostarczania narzędzi analitycznych dla lekarzy.
- Zwiększenia bezpieczeństwa i wykrywania nadużyć.

Platformy e-commerce i duże sklepy internetowe (np. **Allegro**) mogą wykorzystać AI w postaci chatbotów i systemów wykrywania nieuczciwych ofert w celu zwiększenia satysfakcji klientów i obniżenia kosztów obsługi. Może to oznaczać wykorzystanie AI do:

- Zapewnienia lepszego i bardziej spersonalizowanego doświadczenia zakupowego.
- Zwiększenia efektywności wyszukiwania i odkrywania produktów.
- Optymalizacji procesów logistycznych i dostaw.
- Zapewnienia bezpieczeństwa transakcji.
- Usprawnienia obsługi klienta.
- Dostarczania cennych narzędzi i analiz dla sprzedających.
- Ciągłego rozwoju innowacyjnych funkcji i usług.

Instytucje finansowe i banki mogą wykorzystać asystentów AI dla konsultantów oraz do zwiększenia produktywności prac administracyjnych wzrost i skrócenia czasu obsługi. Może to oznaczać wykorzystanie AI do:

- Ulepszenia i personalizacji obsługi klienta poprzez chatboty i wirtualnych asystentów.
- Dostarczania spersonalizowanych ofert produktów i usług.
- Wzmocnienia bezpieczeństwa transakcji i ochrony przed oszustwami finansowymi.
- Automatyzacji procesów wewnętrznych, zwiększając efektywność operacyjną.
- Usprawnienia procesów kredytowych poprzez szybszą i dokładniejszą ocenę ryzyka.
- Wprowadzenia innowacyjnych metod weryfikacji tożsamości, np. wideoweryfikacji.
- Poszukiwania nowych zastosowań AI w bankowości poprzez inicjatywy takie jak hackathony.

Zbliżony zakres działań dotyczy innych przedsiębiorstw obsługujących klientów indywidualnych na skalę masową, takich jak firmy telekomunikacyjne, ubezpieczeniowe, energetyczne, windykacyjne.

Osiągane **korzyści**: redukcja kosztów operacyjnych, zwiększenie przepustowości pracy, minimalizacja pomyłek, odciążenie zespołów od monotonnych obowiązków.

6.2. Poziom 2 – kolejny krok to modyfikacja kluczowych procesów.

Cel: 30–50% poprawy efektywności i skuteczności

Niektóre firmy wykonują kolejny krok, jakim jest głębsza integracja AI w strategiczne procesy biznesowe. Pozwala to na fundamentalną transformację działań operacyjnych.

Przykłady takiej modyfikacji kluczowych procesów w Europie:

- **ASML (Holandia)** – [AI w wykrywaniu defektów i optymalizacji maszyn](#) litograficznych skraca czas rozwoju urządzeń, podnosząc ich precyzję i wydajność.
- **Rolls-Royce (W. Brytania)** – [system predykcyjnej konserwacji silników lotniczych](#) z AI zmniejsza niespodziewane przestoje i obniża koszty serwisu.
- **Ocado (Wlk. Brytania)** – [centra dystrybucji oparte na AI](#) i robotyce (robotyczna kompletacja zamówień sterowania AI) podniosły wydajność kompletacji zamówień i zredukowały koszty.
- **Spotify (Szwecja)** – [algorytmy rekomendacyjne zwiększyły zaangażowanie użytkowników](#) oraz wydłużyły czas korzystania z platformy .

Przykłady takiej modyfikacji kluczowych procesów w Polsce

- Firmy sprzedające produkty konsumenckie z dużym udziałem e-commerce takie jak LPP (Reserved, Cropp, House), Empik. Zastosowanie AI w analizie sprzedaży i trendów modowych pozwala prognozować popyt, redukując nadmiar zapasów i skracając czas reakcji na trendy. Dynamiczne prognozowanie popytu w czasie rzeczywistym przyspieszyło pracę magazynu i obniżyło koszty magazynowania.
- Firmy z branży e-commerce takie jak Allegro tworzą zaawansowane silniki rekomendacyjne i optymalizują logistykę, co zwiększa skuteczność działań operacyjnych.
- Firmy z branży pośrednictwa i rezerwacji takie jak Booksy mogą tworzyć inteligentne algorytmy dopasowujące terminy wizyt, co podnosi liczbę zrealizowanych rezerwacji i redukuje puste terminy.

Korzyści: znacząca redukcja kosztów, szybsze realizacje procesów, poprawa jakości usług, lepsze zarządzanie ryzykiem.

6.3. Poziom 3. Najbardziej zaawansowane wykorzystanie AI to tworzenie nowych produktów i usług

Cel: budowa długoterminowej przewagi konkurencyjnej

Niektóre firmy osiągają najbardziej transformacyjny poziom wykorzystania AI. Polega on na tworzeniu nowych produktów i usług, które bez AI nie mogłyby powstać lub funkcjonować.

Przykłady w Europie i Polsce

- Firma prowadząca badania nad strukturą białek, taka jak **DeepMind** (Wlk. Brytania) – AlphaFold, może wykorzystać AI do przewidywania struktury białek, co może rewolucjonizować farmację i biotechnologię, otwierając nowe ścieżki badań.
- **UiPath (Rumunia)** – RPA wzbogacone o AI umożliwia firmom automatyzację bez głębokiej wiedzy technicznej; wycena firmy przekroczyła 35 mld USD.
- **Infermedica** – system do wstępnej diagnozy medycznej, wspierający lekarzy i pacjentów w priorytetyzacji działań medycznych.
- **ElevenLabs** – generowanie realistycznego głosu syntetycznego do audiobooków, dubbingu i asystentów głosowych.
- **Lekta.ai** – Tworzy wirtualnych agentów do wykorzystania w call centers
- **Solwena** – stworzyła systemy dynamicznego modelowania, prognozowania i zarządzania zużyciem energii w budynkach z użyciem AI.
- **DataWalk** – platforma analityki grafowej do wykrywania nadużyć finansowych i wsparcia wywiadu, używana m.in. przez instytucje rządowe.
- **StethoMe** – inteligentny stetoskop z bazą dźwięków oddechowych; skuteczność wykrywania nieprawidłowości oddechu.
- **SmokedSystems** – dostarcza usługi wczesnego wykrywania pożarów lasów dzięki algorytmom AI wykrywającym dym poprzez analizę obrazu z kamer (computer vision).

Korzyści: stworzenie unikalnej oferty rynkowej, silna pozycja lidera w nowych segmentach, nowe źródła przychodów, wzmacnianie marki innowacyjnej.

Firmy z Europy i Polski, które osiągnęły któryś z trzech poziomów dojrzałości AI – od prostych automatyzacji, przez głęboką transformację procesów, aż po tworzenie zupełnie nowych produktów – osiągają wymierne efekty: podnoszą wydajność, redukują koszty, budują przewagi konkurencyjne i otwierają nowe rynki.

6.4. Przykłady wpływu na wybrane branże AI połączonego z nowoczesnymi urządzeniami. Szanse i zagrożenia

W tej części zawarte są przykłady zmian, które mogą nastąpić w wielu branżach jako konsekwencje adopcji technologii opartych na AI, a także przełomów wynikających z integracji AI z istniejącymi technologiami poza wprost rozumianą domeną produktów i usług cyfrowych. Oznacza to szanse rynkowe, ale również ryzyka, a w każdym razie zmiany, do których dostosować się muszą firmy, których te zmiany dotyczą.

Transport i Logistyka – Pojazdy Autonomiczne

Technologia autonomicznych pojazdów ma potencjał być najbardziej **przełomową** w wielu branżach, fundamentalnie zmieniając istniejące modele biznesowe i tworząc nowe możliwości. Oto niektóre kluczowe sektory, które prawdopodobnie doświadczą najbardziej znaczących zmian:

Przewidywane zmiany

- Autonomiczne ciężarówki i furgonetki mogą działać 365/24/7, obniżając koszty pracy, zwiększając wydajność i potencjalnie przyspieszając czas dostawy. Może to prowadzić do znaczących zmian w transporcie długodystansowym, dostawach „ostatniej mili” i ogólnym zarządzaniu łańcuchem dostaw.
- Autonomiczne pojazdy mogłyby drastycznie obniżyć koszty usług przewozu osób, czyniąc je bardziej dostępnymi i potencjalnie zastępując tradycyjne firmy taksówkarskie, a nawet prywatną własność samochodów dla niektórych.
- Autonomiczne autobusy i wahadłowce mogłyby optymalizować trasy, obniżać koszty operacyjne i poprawiać wydajność oraz dostępność sieci transportu publicznego.
- Autonomiczne statki i samoloty cargo mogłyby prowadzić do bardziej efektywnego transportu towarów, a autonomiczne drony już rewolucjonizują dostawy ostatniej mili i nadzór.
- Autonomiczne pojazdy mogą wykonywać niebezpieczne i powtarzalne zadania w górnictwie i budownictwie, poprawiając bezpieczeństwo i wydajność.
- Autonomiczne ambulanse, wozy strażackie i pojazdy policyjne mogłyby potencjalnie reagować szybciej i skuteczniej na sytuacje awaryjne.
- Ponieważ pasażerowie nie musieliby już skupiać się na prowadzeniu, autonomiczne pojazdy mogłyby stać się mobilnymi centrami rozrywki lub przestrzeniami do pracy.

Szanse

- Działanie 24/7 bez przerw, zwiększona wydajność operacyjna.
- Obniżenie kosztów pracy i transportu.
- Szybsze czasy dostaw.
- Optymalizacja tras i zmniejszenie zużycia paliwa.
- Zwiększona dostępność transportu publicznego i prywatnego.

Zagrożenia i problemy

- Masowa utrata miejsc pracy (kierowcy ciężarówek, taksówkarze, kurierzy).
- Problemy z odpowiedzialnością w przypadku wypadków.
- Zagrożenia cyberbezpieczeństwa i potencjalne awarie systemowe.
- Wysokie koszty początkowe wdrożenia technologii.
- Nerozwiazane kwestie prawne i ubezpieczeniowe.

Transport i Logistyka – Drony Dostawcze

Przewidywane zmiany

- Drony dostawcze mogą zrewolucjonizować dostawy „ostatniej mili”, szczególnie w obszarach miejskich i trudno dostępnych regionach, oferując szybsze i bardziej elastyczne opcje dostawy.
- Firmy kurierskie i e-commerce będą mogły oferować dostawy w ciągu tego samego dnia lub nawet godziny, znacząco zmieniając oczekiwania konsumentów odnośnie do szybkości dostaw.
- Automatyczne systemy zarządzania flotami dronów pozwolą na koordynację setek urządzeń jednocześnie, optymalizując trasy i minimalizując koszty operacyjne.

- Integracja z autonomicznymi pojazdami naziemnymi może stworzyć hybrydowe systemy dostawcze, gdzie pojazdy służą jako mobilne centra dla dronów obsługujących obszar w promieniu kilku kilometrów.
- Autonomiczne mobilne roboty (AMR) z AI nawigują po magazynach, transportują towary i współpracują z pracownikami przy kompletacji zamówień, optymalizując przepływ pracy i redukując błędy.

Szanse

- Szybkie dostawy w trudno dostępnych lokalizacjach
- Obniżenie kosztów dostaw „ostatniej mili”
- Zmniejszenie zatorów drogowych
- Zeroemisyjne opcje transportu (drony elektryczne)

Zagrożenia i problemy

- Problemy z prywatnością i nadzorem
- Zagrożenia dla ruchu lotniczego
- Hałas w obszarach miejskich
- Ograniczenia wynikające z pogody i zasięgu

Roboty Przemysłowe i Coboty

Przewidywane zmiany

- Tradycyjne linie produkcyjne będą ewoluować w kierunku elastycznych, zautomatyzowanych systemów produkcyjnych, gdzie roboty i coboty (roboty współpracujące) pracują obok ludzi.
- Wzrośnie elastyczność produkcji dzięki adaptacyjnym robotom przemysłowym, które mogą szybko dostosowywać się do różnych specyfikacji krótkich serii produktów bez długiego przestoju.
- Mniejsze firmy produkcyjne, które wcześniej nie mogły sobie pozwolić na automatyzację, będą miały dostęp do przystępnych cenowo i łatwych w obsłudze cobotów, demokratyzując robotykę przemysłową.
- Powstaną nowe specjalizacje koncentrujące się na zarządzaniu, nadzorze i konserwacji systemów robotycznych, wymagające kombinacji umiejętności technicznych i zarządczych.
- Fabryki staną się bardziej inteligentne dzięki integracji robotów z systemami IoT i AI, umożliwiając predykcyjne utrzymanie, adaptacyjną produkcję i optymalizację procesów w czasie rzeczywistym.
- Łańcuchy dostaw ulegną transformacji w kierunku większej elastyczności i odporności dzięki zautomatyzowanym systemom magazynowania i kompletacji, które mogą działać 24/7 i szybko dostosowywać się do zmian popytu.

Szanse

- Zwiększenie wydajności produkcji o 30–50%
- Poprawa jakości i powtarzalności produktów
- Wykonywanie niebezpiecznych zadań bez narażania ludzi

- Praca w trudnych warunkach (wysokie temperatury, toksyczne środowiska)
- Optymalizacja procesów dzięki analizie danych w czasie rzeczywistym
- Bezpieczna współpraca z ludźmi w tym samym środowisku pracy
- Elastyczne wsparcie w zmiennych zadaniach produkcyjnych
- Wspomaganie pracowników w ergonomicznie trudnych zadaniach
- Łatwiejsze programowanie i wdrażanie niż w przypadku tradycyjnych robotów

Zagrożenia i problemy

- Redukcja zatrudnienia w przemyśle produkcyjnym
- Wysokie koszty początkowe i utrzymania
- Uzależnienie od dostawców technologii
- Konieczność przekwalifikowania pracowników
- Problemy z akceptacją społeczną i adaptacją pracowników
- Potencjalne zagrożenia bezpieczeństwa przy nieprawidłowej konfiguracji
- Wymagane nowe umiejętności od pracowników

Transport autonomiczny i motoryzacja

Przewidywane zmiany

- Przejście w kierunku autonomicznej i potencjalnie współdzielonej mobilności mogłoby prowadzić do spadku indywidualnej własności samochodów, wpływając na tradycyjne modele sprzedaży producentów samochodów. Mogą oni potrzebować dostosować się, stając się dostawcami usług lub skupiając się na dostarczaniu autonomicznych flot pojazdów.
- Przy oczekiwaniu mniejszej liczby wypadków spowodowanych błędem ludzkim, tradycyjny model ubezpieczeń samochodowych mógłby ulec znacznym zmianom, potencjalnie prowadząc do niższych składek i nowych rodzajów ubezpieczeń skupiających się na ryzykach związanych z technologią i cyberbezpieczeństwem.
- Mniejsza liczba wypadków, zoptymalizowana konstrukcja i eksploatacja pojazdów, mogłyby prowadzić do mniejszego popytu na tradycyjne usługi naprawy samochodów i części zamienne.

Szanse

- Mniejsza liczba wypadków dzięki eliminacji błędu ludzkiego
- Uwolnienie przestrzeni miejskiej zajmowanej przez parkingi
- Zoptymalizowany przepływ ruchu i redukcja korków
- Nowe modele biznesowe oparte na mobilności jako usłudze

Zagrożenia

- Zmiany wartości nieruchomości w zależności od lokalizacji
- Zagrożenie dla tradycyjnego modelu ubezpieczeń samochodowych
- Potrzeba kosztownej modernizacji infrastruktury miejskiej

Opieka zdrowotna – roboty chirurgiczne i asystujące

Przewidywane zmiany

- Zaawansowane roboty chirurgiczne umożliwią przeprowadzanie skomplikowanych zabiegów z dotychczas nieosiągalną precyzją, minimalizując ryzyko błędów ludzkich i skracając czas operacji.
- Powstanie sieci telemedycznych umożliwi chirurgom–ekspertom przeprowadzanie zabiegów zdalnie w odległych lokalizacjach, demokratyzując dostęp do specjalistycznej opieki zdrowotnej.
- Roboty asystujące będą przejmować rutynowe zadania w placówkach medycznych, takie jak transport materiałów, leków i próbek, a także podstawowe pomiary parametrów życiowych pacjentów.
- Autonomiczne roboty do dezynfekcji będą poprawiać standardy sterylności w szpitalach, potencjalnie zmniejszając liczbę zakażeń szpitalnych i związanych z nimi powikłań.
- Egzoszkielety rehabilitacyjne wspomagane AI umożliwią spersonalizowane programy rehabilitacji, adaptujące się w czasie rzeczywistym do postępów pacjenta i jego potrzeb fizjologicznych.

Szanse

- Zwiększona precyzja zabiegów chirurgicznych
- Mniej inwazyjne procedury z krótszym czasem rekonwalescencji
- Dostęp do specjalistycznej opieki w odległych lokalizacjach (telemedycyna)
- Wsparcie w rehabilitacji pacjentów
- Odciążenie personelu medycznego z rutynowych zadań administracyjnych i logistycznych

Zagrożenia

- Wysokie koszty zakupu i utrzymania zaawansowanego sprzętu
- Potencjalne błędy systemu podczas krytycznych procedur
- Problemy z odpowiedzialnością za błędy medyczne
- Dehumanizacja opieki zdrowotnej

Rolnictwo precyzyjne i zrobotyzowane

Przewidywane zmiany

- Autonomiczne traktory, siewniki i kombajny będą operować na polach 24/7, optymalizując czasy zasiewów i zbiorów w oparciu o warunki pogodowe i dojrzałość roślin.
- Drony wyposażone w zaawansowane czujniki będą regularnie monitorować stan upraw, identyfikując wczesne oznaki chorób, niedoborów składników odżywczych czy szkodników, umożliwiając punktową interwencję zamiast ogólnego oprysku.
- Roboty do selektywnego odchwaszczania i precyzyjnego nawożenia zmniejszą zużycie chemikaliów w rolnictwie o 70–90%, minimalizując wpływ na środowisko.
- Systemy hydroponiczne i wertykalnego rolnictwa sterowane przez AI zoptymalizują produkcję żywności w warunkach kontrolowanych, umożliwiając uprawy niezależne od warunków atmosferycznych i bliżej miejsc konsumpcji.

- Inteligentne systemy nawadniania dostosowane do mikroklimatu i rzeczywistych potrzeb roślin zoptymalizują zużycie wody w rolnictwie, które obecnie odpowiada za około 70% globalnego zużycia wody słodkiej.

Szanse

- Zwiększenie wydajności i plonów
- Redukcja zużycia środków ochrony roślin i nawozów
- Precyzyjne zarządzanie uprawami na podstawie analizy danych
- Autonomiczne maszyny rolnicze pracujące 24/7
- Zmniejszenie zależności od ograniczonej siły roboczej

Zagrożenia

- Wysokie koszty początkowe technologii dla małych gospodarstw
- Zwiększenie przepaści technologicznej między dużymi a małymi rolnikami
- Zależność od łączności internetowej i danych
- Potencjalne problemy z cyberbezpieczeństwem systemów rolniczych

Nowe obszary integracji AI ze sprzętem – inteligentne miasta i IoT

Przewidywane zmiany

- Inteligentne systemy zarządzania ruchem będą dynamicznie dostosowywać sygnalizację świetlną w oparciu o rzeczywisty przepływ pojazdów, redukując korki i emisje spalin o 15–30%.
- Sieci czujników miejskich połączone z AI będą monitorować jakość powietrza, poziom hałasu i inne parametry środowiskowe, umożliwiając natychmiastowe reakcje na problemy i lepsze planowanie miejskie.
- Inteligentne systemy zarządzania odpadami zoptymalizują trasy pojazdów odbierających śmieci i zautomatyzują proces sortowania, zwiększając efektywność recyklingu.
- Budynki wyposażone w IoT będą automatycznie dostosowywać zużycie energii do rzeczywistych potrzeb i warunków, potencjalnie redukując zużycie energii o 20–40%.
- Systemy wczesnego ostrzegania o zagrożeniach naturalnych (powódzie, trzęsienia ziemi) będą wykorzystywać dane z czujników i algorytmy predykcyjne do ochrony mieszkańców.

Szanse

- Zoptymalizowane zarządzanie zasobami miejskimi (energia, woda)
- Inteligentne oświetlenie i monitoring środowiskowy
- Poprawa bezpieczeństwa publicznego
- Efektywne zarządzanie odpadami i recyklingiem

Zagrożenia

- Problemy z prywatnością obywateli

- Zagrożenia cyberbezpieczeństwa w infrastrukturze krytycznej
- Wysokie koszty wdrożenia dla mniejszych miast
- Zależność od dostawców technologii

Energetyka i sieci inteligentne

Przewidywane zmiany

- Inteligentne sieci energetyczne (smart grids) będą w czasie rzeczywistym zarządzać produkcją, dystrybucją i konsumpcją energii, dostosowując się do zmiennej podaży energii odnawialnej.
- Systemy magazynowania energii sterowane przez AI zoptymalizują ładowanie i rozładowywanie w oparciu o prognozy produkcji, zapotrzebowania i cen energii.
- Mikrosieci lokalne zwiększą odporność systemu energetycznego na awarie, umożliwiając autonomiczne działanie wydzielonych obszarów podczas problemów z główną siecią.
- Predykcyjne systemy konserwacji infrastruktury energetycznej zidentyfikują potencjalne awarie przed ich wystąpieniem, zmniejszając liczbę przerw w dostawach energii.

Szanse

- Optymalizacja produkcji i dystrybucji energii
- Integracja odnawialnych źródeł energii
- Redukcja strat energii w sieci
- Predykcyjne utrzymanie infrastruktury

Zagrożenia

- Złożoność systemowa zwiększająca ryzyko awarii
- Potencjalne cele dla cyberataków
- Wysokie koszty modernizacji istniejącej infrastruktury

Roboty humanoidalne

Przewidywane zmiany

- Roboty humanoidalne pojawią się jako asystenci w miejscach publicznych, takich jak lotniska, centra handlowe i instytucje rządowe, zapewniając informacje i podstawową pomoc.
- W opiece domowej roboty humanoidalne będą wspierać osoby starsze i niepełnosprawne w codziennych czynnościach, monitorując ich stan zdrowia i zapewniając towarzystwo.
- Roboty ratownicze o budowie humanoidalnej będą operować w niebezpiecznych środowiskach (np. pożary, katastrofy naturalne), docierając do miejsc niedostępnych dla ludzi i wykonując zadania ratunkowe z użyciem narzędzi i pojazdów wcześniej przygotowanych dla człowieka.
- W edukacji roboty humanoidalne będą pełnić rolę spersonalizowanych asystentów nauczyciela, dostosowując metody i tempo nauczania do indywidualnych potrzeb uczniów.
- Eksploracja kosmiczna i głębinowa będzie wykorzystywać roboty humanoidalne jako pierwszych "badaczy", zanim ludzie będą mogli bezpiecznie dotrzeć do nowych środowisk.

Szanse

- Wsparcie osób starszych i niepełnosprawnych w codziennych czynnościach
- Zastąpienie ludzi w zadaniach niebezpiecznych (np. gaszenie pożarów)
- Eksploracja ekstremalnych środowisk (kosmos, głębiny morskie)
- Nowe możliwości w edukacji i terapii

Zagrożenia

- Problemy etyczne i emocjonalne związane z interakcją człowiek-robot
- Potencjalna utrata miejsc pracy w sektorze usługowym
- Wysokie koszty technologii w początkowych fazach wdrożenia
- Niepewność społeczna i regulacyjna

Egzoszkielety wspomagane AI

Przewidywane zmiany

- W przemyśle produkcyjnym i budownictwie egzoszkielety zmniejszą liczbę urazów związanych z pracą fizyczną, jednocześnie zwiększając wydajność pracowników.
- Medyczne egzoszkielety umożliwią osobom z niepełnosprawnościami ruchowymi odzyskanie mobilności, a w niektórych przypadkach nawet powrót do pracy zawodowej.
- Powstanie nowa kategoria zawodów "wspomaganych mechanicznie", gdzie ludzie wyposażeni w egzoszkielety będą wykonywać zadania wymagające zarówno siły fizycznej, jak i ludzkiego osądu.
- Wojskowe zastosowania egzoszkieletów zmienią charakter działań bojowych, umożliwiając żołnierzom przenoszenie cięższego sprzętu na większe odległości i zapewniając lepszą ochronę.
- Zaawansowane egzoszkielety z funkcjami haptycznymi umożliwią operatorom „czucie” środowiska przez roboty, otwierając nowe możliwości w telerobotyce i pracy w skrajnych warunkach.

Szanse

- Zwiększenie siły i wytrzymałości pracowników fizycznych
- Rehabilitacja osób z niepełnosprawnościami ruchowymi
- Redukcja urazów związanych z pracą fizyczną
- Zwiększenie wydajności w budownictwie i produkcji

Zagrożenia

- Wysokie koszty jednostkowe
- Potencjalne problemy zdrowotne związane z długotrwałym użytkowaniem
- Zależność od zasilania i łączności

Systemy produkcji addytywnej (druk 3D) z AI

Przewidywane zmiany

- Systemy druku 3D sterowane przez AI będą automatycznie optymalizować projekty pod kątem wytrzymałości, zużycia materiału i procesu produkcji.
- Wskutek lepszego zaspokojenia potrzeb klientów i spadku kosztów wzrośnie rozmiar rynku druku 3D.
- Personalizacja produktów na masową skalę stanie się ekonomicznie opłacalna, umożliwiając dostosowanie produktów do indywidualnych potrzeb bez znaczącego wzrostu kosztów.

Szanse

- Automatyczna optymalizacja projektów pod kątem wytrzymałości i zużycia materiału
- Tworzenie złożonych struktur niemożliwych do wykonania tradycyjnymi metodami
- Optymalizacja procesu druku uwzględniająca różne parametry fizyczne w różnych miejscach w komorze drukowania

Zagrożenia

- Problemy z prawami własności intelektualnej
- Potencjalne wykorzystanie do tworzenia nielegalnych przedmiotów
- Zmiana struktury zatrudnienia w produkcji tradycyjnej

Edukacja – systemy AI i technologie immersyjne

Przewidywane zmiany

- Spersonalizowane ścieżki edukacyjne dostosowane w czasie rzeczywistym do tempa nauki, stylu poznawczego i zainteresowań każdego ucznia staną się standardem, zastępując ujednolicone programy nauczania.
- Inteligentni tutorzy AI będą wspierać nauczycieli, zapewniając indywidualne wsparcie każdemu uczniowi jednocześnie, monitorując postępy i identyfikując obszary wymagające dodatkowej uwagi.
- Technologie immersyjne (VR/AR) zrewolucjonizują nauczanie przedmiotów wymagających wizualizacji i doświadczenia, umożliwiając wirtualne wycieczki historyczne, eksperymenty naukowe czy eksplorację niedostępnych miejsc.
- Systemy automatycznej oceny będą analizować nie tylko wyniki, ale również proces myślowy ucznia, dostarczając natychmiastowej i konstruktywnej informacji zwrotnej wykraczającej poza tradycyjne ocenianie.
- Edukacja stanie się bardziej ciągłym procesem, wykraczającym poza formalne instytucje edukacyjne, z systemami mikro-certyfikacji potwierdzającymi konkretne umiejętności i kompetencje, uznawanymi przez pracodawców.
- Globalne platformy edukacyjne z automatycznym tłumaczeniem w czasie rzeczywistym demokratyzują dostęp do najlepszych zasobów edukacyjnych niezależnie od lokalizacji geograficznej.
- Modelowanie predykcyjne pomoże wcześniej identyfikować uczniów zagrożonych niepowodzeniem edukacyjnym, umożliwiając proaktywne interwencje zanim problemy się pogłębią.

Szanse

- Demokratyzacja dostępu do wysokiej jakości edukacji niezależnie od lokalizacji geograficznej i statusu społeczno-ekonomicznego
- Rzeczywiste dostosowanie edukacji do indywidualnych potrzeb i możliwości każdego ucznia
- Uwolnienie nauczycieli od rutynowych zadań administracyjnych i oceniania na rzecz mentoringu i wspierania rozwoju kompetencji społecznych
- Zwiększenie zaangażowania uczniów poprzez interaktywne i immersyjne doświadczenia edukacyjne
- Bardziej efektywne nauczanie trudnych i abstrakcyjnych koncepcji dzięki zaawansowanym wizualizacjom i symulacjom
- Wsparcie uczniów ze specjalnymi potrzebami edukacyjnymi poprzez adaptacyjne technologie
- Przygotowanie uczniów do przyszłego rynku pracy poprzez rozwój kompetencji cyfrowych i umiejętności współpracy z systemami AI
- Ułatwienie kształcenia ustawicznego i przekwalifikowania zawodowego odpowiadającego na zmieniające się wymagania rynku pracy

Zagrożenia i problemy

- Pogłębienie nierówności cyfrowych między uczniami z dostępem do zaawansowanych technologii a tymi bez takiego dostępu
- Nadmierne poleganie na systemach AI kosztem rozwoju umiejętności społecznych i współpracy międzyludzkiej
- Problemy z ochroną prywatności uczniów i bezpieczeństwem ich danych edukacyjnych
- Potencjalna redukcja roli nauczycieli i deprecjacja ich autorytetu
- Wysokie koszty wdrożenia i utrzymania zaawansowanych systemów edukacyjnych
- Ryzyko ustandaryzowania myślenia i kreatywności przez algorytmiczne podejście do nauczania
- Uzależnienie od technologii i zmniejszenie odporności systemu edukacji na awarie techniczne
- Trudności w ocenie umiejętności społecznych, etycznych i krytycznego myślenia przez systemy automatyczne
- Wyzwania adaptacyjne dla nauczycieli wymagające fundamentalnej zmiany metodyki nauczania
- Potencjalna redukcja bezpośrednich interakcji społecznych niezbędnych dla prawidłowego rozwoju emocjonalnego

Wpływ AI na wybrane branże – Podsumowanie

Powyżej opisane kierunki zmian oznaczają konieczność zrozumienia szans i zagrożeń, które ze sobą niosą. Członek Rady Nadzorczej powinien być zdolny do strategicznej analizy tych aspektów w odniesieniu do branży, w której działa nadzorowane przedsiębiorstwo.

7. O autorach:

Zapraszamy do kontaktu i komentowania.



Katarzyna Kazior: Menedżer wdrażający AI i innowacje

Pełniła role zarządcze m.in. Prezesa Zarządu Frisco i Dyrektora Transformacji Cyfrowej w Żabka Polska, gdzie odpowiadała za największą transformację cyfrową w Polsce. Współtwórczyni AI Masterclass w ramach CIONET i założycielka startupu opartego na technologii voice AI. Kariere rozwijała także w McKinsey & Company.

Profil LinkedIn: 



Michał Lasocki: Inwestor i menedżer wspierający innowacje

Michał Lasocki od ponad 25 lat działa jako menedżer i inwestor. Jego doświadczenie łączy finanse, technologię i zarządzanie, obejmując venture capital, rynki kapitałowe, transformację energetyczną i sztuczną inteligencję.

Obecnie jako Partner w funduszu inwestycyjnym EEC Ventures wspiera innowacyjne firmy w obszarach transformacji cyfrowej i energetycznej.

Profil LinkedIn: 

8. Linki do stron wykorzystanych w niniejszym dokumencie

https://commission.europa.eu/topics/eu-competitiveness/draghi-report_en

https://www.eib.org/attachments/lucalli/20220211_economic_investment_report_2022_2023_en.pdf

<https://hai.stanford.edu/ai-index>

<https://www.mckinsey.com/featured-insights/artificial-intelligence/tackling-europes-gap-in-digital-and-ai>

<https://www.goldmansachs.com/insights/artificial-intelligence>

<https://flia.org/notice-state-council-issuing-new-generation-artificial-intelligence-development-plan/>

<https://www.goldmansachs.com/insights/articles/what-advanced-ai-means-for-chinas-economic-outlook>

<https://www.goldmansachs.com/insights/articles/chinas-advances-could-boost-ai-impact-on-global-gdp>

https://www.gov.cn/zhengce/content/2015-05/19/content_9784.htm

https://www.uschamber.com/assets/archived/images/final_made_in_china_2025_report_full.pdf

<https://cset.georgetown.edu/publication/notice-of-the-state-council-on-the-publication-of-made-in-china-2025/>

<https://www.aminer.cn/pub/649bb22b98a8018999c4f0b2/china-ai-development-report-2023>

<https://oecd.ai/en/dashboards/policy-areas/PA14>

<https://economicgraph.linkedin.com/content/dam/me/economicgraph/en-us/PDF/AI-in-the-EU-Report.pdf>

<https://www.bcg.com/publications/2025/closing-the-ai-impact-gap>

<https://digital-commons.usnwc.edu/cgi/viewcontent.cgi?article=8417&context=nwc-review>

<https://futureoflife.org/focus-area/artificial-intelligence/>

<https://www.youtube.com/@MinisterstwoCyfryzacji>

<https://www.bankier.pl/wiadomosc/Do-czego-banki-wykorzystuja-AI-Juz-praktycznie-do-wszystkiego-8730952.html>

<https://www.coursera.org/learn/ai-for-everyone>

<https://www.lto.de/recht/hintergruende/h/lg-hamburg-openai-heise-medien-urheberrecht-kuenstliche-intelligenz-ki-scraping-training/>

<https://digital-strategy.ec.europa.eu>

<https://oecd.ai>

<https://artificialintelligenceact.eu>

<https://futureoflife.org>

<https://www.iso.org/committee/6794475.html>

<https://edpb.europa.eu>

<https://www.linkedin.com/pulse/tm-sztuczna-inteligencja-i-projekt-nowelizacji-prawa-zbigniew-oko%C5%84-ol4nf/>

<https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=CELEX:32023R2854>

<https://blog.hightechcampus.com/htce/ai-helps-us-to-squeeze-every-possible-bit-of-performance-out-of-our-machines>.

<https://www.rolls-royce.com/country-sites/india/discover/2018/data-insight-action-latest.aspx>

<https://www.marketingaiinstitute.com/blog/spotify-artificial-intelligence>