

GOTOWOŚĆ TECHNOLOGICZNA I WYKORZYSTANIE SZTUCZNEJ INTELIGENCJI W PRZEDSIĘBIORSTWACH

ANALIZA NA PRZYKŁADZIE
SPECJALNYCH STREF EKONOMICZNYCH

SZYMON
JOPKIEWICZ

**GOTOWOŚĆ TECHNOLOGICZNA
I WYKORZYSTANIE SZTUCZNEJ
INTELIGENCJI
W PRZEDSIĘBIORSTWACH:
ANALIZA NA PRZYKŁADZIE
SPECJALNYCH STREF
EKONOMICZNYCH**

Szymon Jopkiewicz

GOTOWOŚĆ TECHNOLOGICZNA I WYKORZYSTANIE SZTUCZNEJ INTELIGENCJI W PRZEDSIĘBIORSTWACH

ANALIZA NA PRZYKŁADZIE
SPECJALNYCH STREF EKONOMICZNYCH

**SZYMON
JOPKIEWICZ**

ARCHAEGRAPH
Wydawnictwo Naukowe

AUTOR:
SZYMON JOPKIEWICZ

RECENZJA NAUKOWA:
PROF. ZW. DR HAB. PETRO GARASYIM

KOREKTA REDAKTORSKA, SKŁAD I PROJEKT OKŁADKI
KAROL ŁUKOMIAK

© copyright by author & ArchaeGraph

ISBN: 978-83-68410-54-9

**Wersja elektroniczna dostępna
na stronie internetowej wydawcy:**
www.archaeograph.pl

ARCHAEGRAPH
Wydawnictwo Naukowe

ŁÓDŹ 2025

SPIS TREŚCI

WSTĘP	7
KOMPETENCJE INFRASTRUKTURALNE ORGANIZACJI W ZAKRESIE SZTUCZNEJ INTELIGENCJI	11
USŁUGI ZARZĄDZANIA DANYMI I ARCHITEKTURA AI	11
SIECI, CHMURA OBLICZENIOWA I ZASOBY OBLICZENIOWE	22
BEZPIECZEŃSTWO DANYCH JAKO FUNDAMENT INFRASTRUKTURY AI	31
STRATEGICZNE PODEJŚCIE ORGANIZACJI DO SZTUCZNEJ INTELIGENCJI	42
INTEGRACJA AI Z PLANOWANIEM BIZNESOWYM I STRATEGICZNYM	42
ROLA KIEROWNICTWA I ZARZĄDÓW W INWESTYCJACH W AI	49
PROAKTYWNE PODEJŚCIE: EKSPERYMENTOWANIE, KULTURA INNOWACJI I POSZUKIWANIE NOWYCH ZASTOSOWAŃ AI	58
ZASTOSOWANIE SZTUCZNEJ INTELIGENCJI W MARKETINGU	66
OD SIM DO INTELIGENTNYCH EKOSYSTEMÓW DANYCH AI W MARKETINGU	66
PLANOWANIE MARKETINGOWE WSPIERANE PRZEZ ALGORYTMY PREDYKCYJNE I PERSONALIZACJĘ	78
IMPLEMENTACJA DZIAŁAŃ MARKETINGOWYCH I PRZEWAGA KONKURENCYJNA PRZEDSIĘBIORSTW	89

BADANIA EMPIRYCZNE NAD WDRAŻANIEM AI W PRZEDSIĘBIORSTWACH	102
METODOLOGIA BADAŃ ANKIETOWYCH I CHARAKTERYSTYKA PRÓBY	102
ANALIZA WYNIKÓW BADAŃ (INFRASTRUKTURA, STRATEGIA, MARKETING, WYNIKI ORGANIZACYJNE)	106
WERYFIKACJA HIPOTEZ BADAWCZYCH	117
DYSKUSJA WYNIKÓW I IMPLIKACJE DLA TEORII ORAZ PRAKTYKI W PRZEDSIĘBIORSTWACH SSE	123
ZAKOŃCZENIE	126
BIBLIOGRAFIA	129

Wstęp

Współczesna gospodarka przechodzi transformację, której siłą napędową są technologie cyfrowe. Wśród nich szczególne miejsce zajmuje sztuczna inteligencja, która ewoluowała w kierunku fundamentalnego narzędzia kształtującego innowacyjność, efektywność operacyjną i, co najważniejsze, strategiczną przewagę konkurencyjną przedsiębiorstw. AI, definiowana jako zbiór technologii zdolnych do percepcji, rozumowania, uczenia się i autonomicznego działania, rewolucjonizuje modele biznesowe, wymuszając na organizacjach fundamentalną rewizję dotychczasowych paradygmatów zarządzania. Przejście od tradycyjnych, reaktywnych procesów decyzyjnych do proaktywnych, opartych na danych ekosystemów staje się warunkiem sine qua non przetrwania i rozwoju w coraz bardziej dynamicznym otoczeniu.

Kluczowym zasobem strategicznym w tej nowej rzeczywistości nie jest już samo posiadanie technologii, lecz zdolność organizacji do szybszego niż konkurenci uczenia się, adaptacji i rekonfiguracji zasobów. Zdolność ta, określana mianem przewagi kognitywnej, opiera się na ciągłej analizie danych w czasie rzeczywistym i staje się najcenniejszym, najtrudniejszym do skopiowania aktywem. Jednak pomimo rosnącej świadomości potencjału AI, jej praktyczna implementacja wciąż napotyka na liczne wyzwania. Zjawisko to jest szczególnie widoczne w Polsce, gdzie obserwujemy krytyczną lukę wdrożeniową. Z jednej strony polskie środowisko naukowe może poszczycić się znaczącym dorobkiem w dziedzinie sztucznej inteligencji, co plasuje kraj w czołówce Unii Europejskiej pod względem liczby publikacji. Z drugiej strony, poziom adopcji rozwiązań AI w polskich przedsiębiorstwach należy do najniższych w Europie, a firmy wciąż z oporem wdrażają technologie chmurowe czy zaawansowane platformy analityczne. Rozbieżność ta dowodzi, że bariera nie leży w braku dostępu do technologii, ale w deficycie dojrzałości organizacyjnej, strategicznej wizji, kompetencji cyfrowych oraz kultury sprzyjającej innowacjom.

W niniejszej pracy podjęto próbę kompleksowej analizy tego fenomenu w unikalnym kontekście Specjalnych Stref Ekonomicznych (SSE). Z założenia mające pełnić rolę inkubatorów innowacji i konkurencyjności, SSE stanowią fascynujące pole badawcze. Z jednej strony, oferują one sprzyjający ekosystem – ulgi podatkowe, nowoczesną infrastrukturę i efekt aglomeracji, które obniżają bariery inwestycyjne. Z drugiej strony, to samo środowisko generuje specyficzne wyzwania, takie jak zaostzona konkurencja, niedobory wykwalifikowanej kadry czy trudności w absorpcji zaawansowanej wiedzy od globalnych liderów. Analiza gotowości technologicznej i wykorzystania AI w przedsiębiorstwach działających w SSE pozwala nie tylko na precyzyjną diagnozę czynników sukcesu i porażki w procesie transformacji cyfrowej, ale także na sformułowanie rekomendacji o znaczeniu dla całej gospodarki.

Głównym celem poznawczym i aplikacyjnym pracy jest zidentyfikowanie, a następnie empiryczna weryfikacja kluczowych czynników organizacyjnych i technologicznych wpływających na poziom wdrożenia sztucznej inteligencji oraz jej związek z wynikami biznesowymi przedsiębiorstw. Na podstawie przeglądu literatury przedmiotu oraz analizy danych empirycznych zebranych w ramach badania ankietowego, w pracy postawiono i poddano weryfikacji następujące hipotezy badawcze:

- H1: Poziom wdrożenia sztucznej inteligencji (mierzony indeksem AI) jest skorelowany z wynikami biznesowymi przedsiębiorstwa, w tym z satysfakcją klientów, rentownością oraz innowacyjnością.
- H2: Przedsiębiorstwa z branży technologicznej charakteryzują się wyższym poziomem zaawansowania wdrożeń AI niż przedsiębiorstwa działające w innych sektorach gospodarki.
- H3: Stosowanie wybranych narzędzi AI (na przykładzie platformy Mailchimp) w działaniach marketingowych w istotny sposób oddziałuje na wyniki marketingowe oraz ogólną rentowność firm.
- H4: Wielkość przedsiębiorstwa pełni rolę moderatora w relacji pomiędzy stopniem wdrożenia AI a osiąganymi wynikami biznesowymi, co oznacza, że wpływ AI na wyniki jest różny w zależności od wielkości firmy.
- H5: Występuje związek między stosowaniem zaawansowanych praktyk marketingowych (takich jak dynamiczna alokacja zasobów czy monitorowanie potrzeb klientów w czasie rzeczywistym) a wynikami biznesowymi przedsiębiorstw.

Struktura pracy została podporządkowana logice przyjętego procesu badawczego, prowadząc czytelnika od fundamentów teoretycznych, przez szczegółową analizę zastosowań, aż do empirycznej weryfikacji postawionych hipotez.

Rozdział pierwszy, zatytułowany „Kompetencje infrastrukturalne organizacji w zakresie sztucznej inteligencji”, stanowi dogłębną analizę technologicznych fundamentów gotowości na wdrożenie AI. Podkreślono w nim, że skuteczna implementacja sztucznej inteligencji wymaga synergii między dojrzałą strategią zarządzania danymi a elastyczną, skalowalną architekturą systemową. W rozdziale prześledzono ewolucję architektur danych – od tradycyjnych hurtowni, przez koncepcję jeziora danych (Data Lake), aż po nowoczesny, hybrydowy model Data Lakehouse, który łączy elastyczność z mechanizmami ładu korporacyjnego. Szczegółowo omówiono takie zagadnienia jak potoki przetwarzania danych (ETL), inżynieria cech oraz fundamentalna rola ładu danych w zapewnieniu jakości, bezpieczeństwa i etyki. Dalsza część rozdziału poświęcona jest analizie zasobów obliczeniowych (GPU, TPU) oraz roli chmury obliczeniowej i platform AI-as-a-Service, które demokratyzują dostęp do zaawansowanych technologii. Rozdział zamyka omówienie nowoczesnych architektur sieciowych (5G, Edge Computing).

Rozdział drugi, „Strategiczne podejście organizacji do sztucznej inteligencji”, przenosi akcent z technologii na wymiar strategiczny, kulturowy i przywódczy. Argumentuje się w nim, że prawdziwa wartość AI nie tkwi w optymalizacji istniejących procesów, lecz w redefinicji modeli biznesowych. Rozdział identyfikuje kluczowe bariery hamujące ten proces – od twardych (finansowych, kompetencyjnych) po miękkie (strategiczne, kulturowe). Szczególną uwagę poświęcono roli kadry zarządzającej i rad nadzorczych, które odpowiadają za alokację kapitału, zarządzanie ryzykiem i ustanowienie ram ładu korporacyjnego dla AI. Przedstawiono ewolucję modeli przywództwa w kierunku „AI Leadership”, wymagającego hybrydowych kompetencji technicznych, etycznych i transformacyjnych. Podkreślono również znaczenie proaktywnego podejścia do innowacji poprzez metodyczne eksperymentowanie (projekty PoC) i budowanie kultury organizacyjnej opartej na zaufaniu i ciągłym uczeniu się.

Rozdział trzeci, „Zastosowanie sztucznej inteligencji w marketingu”, stanowi studium przypadku ilustrujące głęboką transformację jednego z kluczowych obszarów funkcjonalnych firmy. Śledzi on ewolucję od tradycyjnych Systemów Informacji Marketingowej (SIM) do inteligentnych, autonomicznych ekosystemów danych napędzanych przez AI. Wykorzystując teoretyczną ramę czterech poziomów inteligencji AI (mechanicznej, analitycznej, intuicyjnej i empatycznej),

szczegółowo opisano, jak sztuczna inteligencja wspiera każdy etap zarządzania informacją – od gromadzenia i analizy, przez generowanie wniosków, aż po automatyzację i hiperpersonalizację działań. W rozdziale pokazano, jak AI, a w szczególności technologie generatywne (GenAI), redefiniują ścieżkę klienta i przesuwają źródło przewagi konkurencyjnej z samego produktu na całościowe, spersonalizowane doświadczenie klienta (CX).

Rozdział czwarty, „Badania empiryczne nad wdrażaniem AI w przedsiębiorstwach”, stanowi empiryczny element pracy. W rozdziale zaprezentowano szczegółową metodologię przeprowadzonych badań ankietowych, charakterystykę próby badawczej (177 przedsiębiorstw, w tym z SSE) oraz wyniki analiz statystycznych. Ta część pracy koncentruje się na formalnej weryfikacji pięciu hipotez badawczych sformułowanych na początku pracy. Przy użyciu metod statystycznych, zbadano relacje między poziomem wdrożenia AI a wynikami biznesowymi, wpływ branży i wielkości firmy oraz znaczenie zaawansowanych praktyk marketingowych. Pracę zamyka dyskusja uzyskanych wyników w kontekście istniejącej literatury oraz sformułowanie implikacji teoretycznych i praktycznych dla menedżerów.

Kompetencje infrastrukturalne organizacji w zakresie sztucznej inteligencji

Usługi zarządzania danymi i architektura AI

Gotowość technologiczna przedsiębiorstwa do wdrożenia i efektywnego wykorzystania sztucznej inteligencji (AI) nie jest prostą sumą odizolowanych komponentów technologicznych. Stanowi ona wynik głębokiej, synergicznej relacji pomiędzy dojrzałą strategią zarządzania danymi a świadomie zaprojektowaną, elastyczną architekturą systemową. W tym złożonym ekosystemie dane pełnią rolę fundamentalnego surowca, bez którego nawet najbardziej zaawansowane algorytmy pozostają bezużyteczne. Z kolei architektura AI jest zaawansowaną „rafinerią”, która przekształca ten surowiec w wartość biznesową – predykcje, automatyzację i nowe modele operacyjne. Bez efektywnego zarządzania cyklem życia danych, architektura nie ma czego przetwarzać; bez odpowiedniej architektury, potencjał danych pozostaje niewykorzystany.

Ewolucja systemów sztucznej inteligencji wymusiła fundamentalną zmianę w paradygmatach przechowywania i zarządzania danymi. Tradycyjne hurtownie danych, zoptymalizowane pod kątem analizy danych ustrukturyzowanych i zorientowane na zapytania SQL, okazały się niewystarczające w obliczu wymagań nowoczesnych modeli AI. Systemy te, zwłaszcza z obszaru uczenia głębokiego, wymagają dostępu do ogromnych wolumenów danych nieustrukturyzowanych i częściowo ustrukturyzowanych, takich jak tekst, obrazy, pliki wideo, logi telemetryczne czy dane z mediów społecznościowych (Hahn, 2025; Tewari, 2025). W odpowiedzi na te wyzwania wyłoniła się koncepcja jeziora danych (Data Lake) jako centralnego, skalowalnego repozytorium, zdolnego pomieścić wszystkie dane przedsiębiorstwa w ich surowej, natywnej formie.

Kluczowe cechy, które definiują jezioro danych i czynią je fundamentem nowoczesnej infrastruktury AI, to przede wszystkim elastyczność, efektywność kosztowa i gotowość do zaawansowanej analityki. W przeciwieństwie do hurtowni, jeziora danych nie narzucają schematu w momencie zapisu danych, lecz pozwalają na jego elastyczne definiowanie w momencie odczytu. Umożliwia to przechowywanie różnorodnych typów danych bez konieczności ich wcześniejszej transformacji (Samola, 2025). Co więcej, architektura jezior danych opiera się na oddzieleniu warstwy obliczeniowej od warstwy przechowywania, co pozwala na niezależne skalowanie obu komponentów i wykorzystanie różnych silników analitycznych (np. Spark, Presto) do przetwarzania tych samych danych. Z perspektywy kosztowej, jeziora danych wykorzystują tanie, chmurowe magazyny obiektowe, takie jak Amazon S3, Azure Blob Storage czy Google Cloud Storage, które oferują niemal nieograniczoną skalowalność, wysoką trwałość i niższy koszt przechowywania w porównaniu z tradycyjnymi bazami danych (Orr, 2025).

Jednakże, elastyczność jezior danych niesie ze sobą istotne ryzyko. Bez rygorystycznego ładu informacyjnego, zarządzania metadanymi i kontroli jakości, jezioro danych może szybko przekształcić się w tzw. „bagno danych” – chaotyczny, niezorganizowany i niespójny zbiór, w którym odnalezienie wartościowych informacji staje się niemożliwe, a zaufanie do danych ulega erozji (Amola, 2025). Dlatego też kluczowe staje się wdrożenie zautomatyzowanych procesów zarządzania metadanymi, często wspieranych przez algorytmy AI, które potrafią automatycznie tagować, klasyfikować i katalogować napływające dane.

W odpowiedzi na dylemat wyboru między ustrukturyzowanymi, ale mało elastycznymi hurtowniami danych a elastycznymi, lecz trudnymi w zarządzaniu jeziorami danych, wyłonił się nowy paradygmat architektoniczny: Data Lakehouse. Po raz pierwszy opisana w 2020 roku, architektura ta stanowi hybrydę, która łączy w sobie najlepsze cechy obu poprzednich generacji: niskokosztowe, skalowalne przechowywanie danych w otwartych formatach, typowe dla jezior danych, z solidnymi mechanizmami zarządzania, transakcyjnością (ACID) i optymalizacją zapytań, znanymi z hurtowni danych (Janssen, 2022). Celem Lakehouse jest stworzenie jednego, zunifikowanego systemu, który może obsługiwać zarówno tradycyjną analitykę biznesową (BI), jak i zaawansowane zastosowania z obszaru data science i uczenia maszynowego, eliminując potrzebę utrzymywania dwóch oddzielnych, często niespójnych systemów (Armbrust et al., 2021).

Fundamentem architektury Lakehouse jest warstwa taniego, chmurowego magazynu obiektowego, na której przechowywane są dane w otwartych formatach plikowych, takich jak Apache Parquet. Kluczowym elementem, który odróżnia

Lakehouse od zwykłego jeziora danych, jest dodanie warstwy metadanych transakcyjnych, zaimplementowanej za pomocą otwartych formatów tabelarycznych, takich jak Apache Iceberg, Apache Hudi czy Delta Lake. Ta warstwa umożliwia realizację operacji typowych dla baz danych, takich jak transakcje ACID, ewolucja schematu, a także „podróże w czasie”, czyli możliwość odpytywania historycznych wersji danych, co jest kluczowe dla audytowalności i powtarzalności eksperymentów AI (Dubey, 2025). Dzięki temu, analitycy mogą wykonywać zapytania SQL bezpośrednio na danych w jeziorze z wydajnością zbliżoną do hurtowni, podczas gdy data scientists mogą pracować na tych samych, spójnych danych, używając narzędzi takich jak Apache Spark do trenowania modeli. Architektura Lakehouse jest zazwyczaj pięciowarstwowa i obejmuje warstwy: pozyskiwania, przechowywania, metadanych, API oraz konsumpcji. Wdrożenie tego paradygmatu pozwala organizacjom na redukcję redundancji danych, obniżenie kosztów, poprawę ładu informacyjnego i, co najważniejsze, na stworzenie jednolitego źródła prawdy dla wszystkich rodzajów obciążeń analitycznych (Jonker & Gomstyn, 2024).

Przekształcenie surowych, często chaotycznych danych z jeziora danych w wysokiej jakości, spójne zbiory treningowe, gotowe do użycia przez modele uczenia maszynowego, odbywa się w ramach uporządkowanego procesu zwanego potokiem przetwarzania danych. Jest to najbardziej czasochłonny i pracochłonny, a jednocześnie najważniejszy etap w cyklu życia projektu AI, którego jakość determinuje końcowy sukces modelu (Hochkamp et al., 2025). Działanie tego procesu opiera się na fundamentalnej zasadzie „śmieci na wejściu, śmieci na wyjściu”, która podkreśla, że nawet najbardziej zaawansowany algorytm nie wygeneruje wartościowych wyników na podstawie błędnych, niekompletnych lub zaszumionych danych (Ndung'u, 2022).

Typowy potok przetwarzania danych dla uczenia maszynowego składa się z kilku kluczowych etapów:

1. Pozyskiwanie - proces gromadzenia danych z różnorodnych źródeł, takich jak bazy danych, systemy ERP, urządzenia IoT czy zewnętrzne API. Współczesne architektury zmierzają w kierunku w pełni zarządzanych, sterowanych zdarzeniami systemów pozyskiwania danych, które wspierają mechanizmy Change Data Capture (CDC). Pozwala to na niemal natychmiastowe przechwytywanie zmian w danych źródłowych i umożliwia analizę w czasie rzeczywistym, co jest kluczowe dla dynamicznych zastosowań AI. Centralizacja tego etapu w ramach jeziora danych

- zapobiega duplikacji wysiłków różnych zespołów i zapewnia spójność danych w całej organizacji (Longo & Saad, 2025);
2. Czyszczenie i wstępne przetwarzanie- na tym etapie surowe dane poddawane są rygorystycznej kontroli jakości. Obejmuje to identyfikację i korygowanie błędów, usuwanie zduplikowanych rekordów, obsługę brakujących wartości (np. poprzez imputację, czyli zastępowanie ich wartościami szacowanymi) oraz wykrywanie i neutralizację wartości odstających, które mogłyby zakłócić proces uczenia modelu (Guo et al., 2023; Martins et al., 2025);
 3. Transformacja i inżynieria cech - jest to proces kreatywnego przekształcania danych w celu wydobycia z nich cech (zmiennych), które najlepiej opisują badany problem i mogą być efektywnie wykorzystane przez algorytm. Techniki stosowane na tym etapie to m.in. normalizacja (skalowanie wartości do wspólnego zakresu), dyskretyzacja (zamiana zmiennych ciągłych na kategorie) oraz tworzenie nowych, złożonych cech na podstawie istniejących (np. obliczanie wieku klienta na podstawie daty urodzenia) (Mumuni, 2024; Lopez-Miguel, 2021);
 4. Wzbogacanie i etykietowanie - dane są wzbogacane o dodatkowy kontekst, np. poprzez dołączenie danych z zewnętrznych źródeł lub dodanie metadanych. W przypadku uczenia nadzorowanego, kluczowym krokiem jest etykietowanie, czyli przypisywanie każdej próbce danych poprawnej odpowiedzi (etykiety), której model ma się nauczyć. Jest to często proces manualny lub półautomatyczny, wymagający dużej precyzji (Hahn, 2025; Dilmegani, 2025).

Złożoność i wieloetapowość tego procesu sprawiają, że jego automatyzacja staje się priorytetem. Pojawiają się zaawansowane techniki, takie jak ML-driven data preparation, w których algorytmy uczenia maszynowego, np. algorytmy ewolucyjne, są wykorzystywane do automatycznego projektowania i optymalizacji całego potoku przetwarzania danych, identyfikując optymalną sekwencję kroków dla danego problemu (Krajnc et al., 2022).

W nawiązaniu do powyższego na uwagę zasługuje ład danych oznaczający strategiczne ramy obejmujące zbiór procesów, ról, polityk, standardów i metryk, które zapewniają efektywne i bezpieczne zarządzanie danymi jako kluczowym zasobem przedsiębiorstwa. W kontekście sztucznej inteligencji jego rola jest nie do przecenienia, ponieważ jakość, etyka, transparentność i zgodność z prawem systemów AI są bezpośrednim odzwierciedleniem ład danych, na których zostały

one wytrenowane (Tewari, 2025; Quinnox, 2025). Bez solidnych fundamentów w postaci ładu danych, organizacje ryzykują budowanie potężnych systemów AI na niestabilnym gruncie, co prowadzi do błędnych decyzji, naruszeń regulacyjnych i utraty zaufania interesariuszy.

Fundamentem skutecznego ładu danych jest zapewnienie ich wysokiej jakości, ocenianej w kilku kluczowych wymiarach. Należą do nich: dokładność, czyli zgodność danych z rzeczywistością; spójność, oznaczająca brak sprzeczności w danych w różnych systemach; kompletność, czyli brak brakujących wartości; aktualność, zapewniająca, że dane odzwierciedlają bieżący stan; oraz adekwatność, gwarantująca, że dane są odpowiednie do zamierzonego celu (Dilmegani, 2025). Niska jakość w którymkolwiek z tych wymiarów nieuchronnie prowadzi do degradacji wydajności modelu AI.

W erze AI tradycyjne podejście do ładu danych musi zostać rozszerzone o nowe komponenty, specyficznie adresujące wyzwania związane z uczeniem maszynowym:

- Pochodzenie danych - jest to zdolność do precyzyjnego śledzenia całego cyklu życia danych – od ich pierwotnego źródła, przez wszystkie etapy transformacji w potoku danych, aż po ich wykorzystanie w procesie trenowania i wnioskowania konkretnej wersji modelu AI. Pełna widoczność pochodzenia danych jest niezbędna dla zapewnienia audytowalności, transparentności („wyjaśnialności” decyzji AI) oraz dla szybkiego diagnozowania problemów, gdy wydajność modelu spada (Bomma, 2024);
- Zarządzanie metadanymi - w kontekście AI, zarządzanie metadanymi wykracza poza proste katalogowanie. Obejmuje ono automatyczne tagowanie i klasyfikację danych w celu identyfikacji informacji wrażliwych (np. danych osobowych), co jest kluczowe dla zapewnienia zgodności z regulacjami takimi jak RODO. Metadane dostarczają również kontekstu niezbędnego do zrozumienia danych i ich prawidłowego wykorzystania (Samola, 2025);
- Zarządzanie ryzykiem i zgodnością - nowe regulacje, takie jak unijny AI Act, wprowadzają podejście oparte na ryzyku, klasyfikując systemy AI i nakładając na ich twórców i użytkowników odpowiednie obowiązki. Skuteczny ład danych musi proaktywnie integrować te wymagania, zapewniając, że dane używane do trenowania, walidacji i wnioskowania są zgodne z obowiązującym prawem na każdym etapie cyklu życia (Chojnowski, 2025);

- Etyka i zarządzanie uprzedzeniami - modele AI uczą się na podstawie danych historycznych, które mogą zawierać ukryte uprzedzenia (np. dotyczące płci, rasy czy statusu społeczno-ekonomicznego). Jeśli nie zostaną one zidentyfikowane i zneutralizowane, AI nie tylko je powieli, ale często również wzmocni. Dlatego ład danych musi obejmować mechanizmy do aktywnego wykrywania i mitygacji uprzedzeń w zbiorach treningowych, np. poprzez techniki reweighting lub resampling, oraz regularne audyty fairness (Quinnox, 2025).

Wdrożenie tych komponentów prowadzi do powstania dynamicznej pętli sprzężenia zwrotnego, która łączy ład danych z operacjonalizacją modeli AI (MLOps). W tym modelu, procesy MLOps, takie jak ciągłe monitorowanie wydajności modelu w środowisku produkcyjnym, dostarczają kluczowych informacji zwrotnych dla systemu ład danych. Gdy system monitorujący wykryje spadek dokładności modelu lub pojawienie się tendencyjnych predykcji, uruchamiana jest procedura diagnostyczna. Wykorzystując narzędzia do śledzenia pochodzenia danych, analitycy mogą cofnąć się wzdłuż potoku danych, aby zidentyfikować źródło problemu – czy to nową, niskiej jakości partię danych, czy błąd w transformacji. Na podstawie tej diagnozy, polityki ład danych są aktualizowane, np. poprzez wprowadzenie nowych reguł walidacji danych lub modyfikację procesu ich pozyskiwania. W ten sposób ład danych przekształca się z reaktywnego, skoncentrowanego na zgodności z regulacjami centrum kosztów, w proaktywny, napędzany danymi silnik ciągłego doskonalenia systemów AI. Zdolność organizacji do wdrożenia i efektywnego zarządzania tą pętlą sprzężenia zwrotnego staje się miarą jej dojrzałości w zakresie AI i kluczową kompetencją infrastrukturalną.

Platformy chmury publicznej, takie jak Amazon Web Services (AWS), Microsoft Azure i Google Cloud Platform (GCP), stały się de facto standardem i podstawową infrastrukturą dla rozwoju i wdrażania sztucznej inteligencji w przedsiębiorstwach. Oferują one bezprecedensową skalowalność, elastyczność i model płatności za faktyczne użycie, eliminując potrzebę kosztownych, początkowych inwestycji w fizyczną infrastrukturę (Luitse, 2024). Ta zmiana paradygmatu znacząco obniżyła barierę wejścia, demokratyzując dostęp do zaawansowanych technologii AI dla szerokiego spektrum firm. Kluczową rolę w tej demokratyzacji odgrywają w pełni zarządzane platformy AI/ML, takie jak AWS SageMaker, Google Vertex AI i Azure Machine Learning. Usługi te abstrahują znaczną część złożoności związanej z zarządzaniem infrastrukturą, dostarczając zintegrowane środowiska, które obejmują cały cykl życia modelu uczenia maszynowego –

od przygotowania danych, przez trenowanie i strojenie hiperparametrów, po wdrożenie, monitorowanie i zarządzanie (Oladele, 2025). Dzięki temu zespoły data science mogą skupić się na budowie modeli i tworzeniu wartości biznesowej, zamiast na konfiguracji serwerów i zarządzaniu zależnościami oprogramowania. Dostawcy chmury oferują kompleksowe ekosystemy usług, które wspierają każdy aspekt projektu AI. Obejmują one skalowalne magazyny danych (np. Amazon S3, Google BigQuery), elastyczną moc obliczeniową (np. maszyny wirtualne Amazon EC2, konteneryzowane aplikacje w Google Cloud Run) oraz szeroką gamę gotowych do użycia, pre-trenowanych modeli AI w formie API (np. Azure Cognitive Services do rozpoznawania mowy i obrazu, Google Vision API) (Oladele, 2025). W rezultacie, wybór dostawcy chmury staje się strategiczną decyzją, która głęboko wpływa na cały stos technologiczny AI w przedsiębiorstwie, determinując dostępne narzędzia, architekturę i potencjalne koszty (Luitse, 2024).

Podstawą wydajności systemów AI, zwłaszcza modeli głębokiego uczenia, jest wyspecjalizowana infrastruktura sprzętowa zdolna do wykonywania ogromnej liczby operacji matematycznych w sposób równoległy. Dwa typy procesorów zdominowały ten obszar: procesory graficzne (GPU) i tensorowe jednostki przetwarzające (TPU).

Procesory graficzne (GPU), pierwotnie zaprojektowane do renderowania grafiki komputerowej, okazały się rewolucyjne dla AI. Ich architektura, składająca się z tysięcy mniejszych rdzeni obliczeniowych, jest idealnie przystosowana do wykonywania operacji na macierzach i wektorach, które stanowią fundament działania sieci neuronowych. Siłą GPU jest ich wszechstronność i elastyczność; mogą być używane do szerokiego spektrum zadań obliczeniowych, nie tylko związanych z AI. Firma NVIDIA, dzięki swojej platformie programistycznej CUDA, zdominowała rynek, posiadając około 80% udziału w segmencie akceleratorów AI (ByteBridge, 2025; Dąbrowski, 2023).

Tensorowe jednostki przetwarzające to specjalistyczne układy scalone (ASIC) zaprojektowane od podstaw przez Google w jednym celu: maksymalnego przyspieszenia operacji na tensorach, które są kluczowymi strukturami danych w uczeniu głębokim. W przeciwieństwie do GPU, TPU nie są uniwersalne. Ich architektura jest zoptymalizowana pod kątem konkretnych zadań AI, co pozwala im osiągnąć znacznie wyższą wydajność i efektywność energetyczną w tych zastosowaniach (Jouppi et al., 2017; Kasinadhsarma, 2024).

Analiza porównawcza obu technologii ujawnia strategiczny kompromis. TPU wykazują od 2 do 3 razy lepszą wydajność na wat w porównaniu do GPU i mogą być od 15 do 30 razy szybsze w trenowaniu specyficznych modeli sieci

neuronowych. Oznacza to, że dla dużych, zunifikowanych obciążeń AI, takich jak trenowanie modeli językowych na masową skalę, TPU oferują znacznie niższy całkowity koszt posiadania (TCO), zwłaszcza gdy są wykorzystywane w modelu chmurowym. Z drugiej strony, GPU oferują większą elastyczność programistyczną, szersze wsparcie w różnych frameworkach i niższy próg wejścia dla mniejszych, bardziej zróżnicowanych projektów. Wybór między GPU a TPU nie jest więc kwestią techniczną, lecz strategiczną decyzją, która musi uwzględniać skalę operacji, rodzaj obciążeń i długoterminowe cele przedsiębiorstwa (ByteBridge, 2025).

Wybór odpowiedniej architektury systemowej jest jedną z najważniejszych decyzji strategicznych w procesie wdrażania AI. Determinuje ona nie tylko wydajność i skalowalność, ale również aspekty takie jak koszty, bezpieczeństwo, prywatność danych i zdolność do działania w czasie rzeczywistym. Współczesne przedsiębiorstwa mają do dyspozycji trzy główne paradygmaty architektoniczne: scentralizowany, sfederowany i brzegowy, a także coraz częściej stosowane modele hybrydowe:

- Architektura scentralizowana - jest to dominujący i najbardziej dojrzały model, w którym dane z różnych źródeł są gromadzone, przechowywane i przetwarzane w centralnym repozytorium, najczęściej zlokalizowanym w chmurze publicznej lub prywatnej. Ten paradygmat pozwala na wykorzystanie ogromnej, niemal nieograniczonej mocy obliczeniowej chmury do trenowania bardzo złożonych modeli, takich jak duże modele językowe (LLM), oraz do prowadzenia analiz na zbiorach danych typu Big Data. Główne zalety tego podejścia to wysoka skalowalność, łatwość zarządzania zasobami i dostęp do zaawansowanych, zintegrowanych platform AI/ML. Jednakże, model ten ma również istotne wady: wysoką latencję związaną z koniecznością przesyłania danych do i z chmury, znaczne koszty transferu danych oraz fundamentalne wyzwania związane z prywatnością i suwerennością danych (Yao et al., 2021);
- Architektura zdecentralizowana - zamiast gromadzić dane w jednym miejscu, model AI jest wysyłany do zdecentralizowanych urządzeń lub lokalizacji (klientów), gdzie jest trenowany na lokalnych danych. Następnie, do centralnego serwera odsyłane są jedynie zaktualizowane parametry (wagi) modelu, a nie surowe dane. Serwer agreguje te parametry w celu stworzenia ulepszanego, globalnego modelu, który jest ponownie dystrybuowany do klientów. Główną zaletą tego podejścia jest

fundamentalne wzmocnienie prywatności i bezpieczeństwa – wrażliwe dane nigdy nie opuszczają lokalizacji klienta. Czyni to uczenie sfederowane idealnym rozwiązaniem dla sektorów silnie regulowanych, takich jak opieka zdrowotna czy finanse, oraz dla scenariuszy współpracy między organizacjami, które nie mogą lub nie chcą dzielić się swoimi danymi. Wyzwania obejmują złożoność zarządzania komunikacją, radzenie sobie z heterogenicznością danych i zasobów obliczeniowych u klientów oraz zabezpieczenie procesu agregacji przed potencjalnymi atakami (Rahman, 2025; Kairouz et al., 2021);

- Architektura brzegowa (Edge AI) - ten paradygmat polega na przeniesieniu obliczeń AI (zwłaszcza etapu wnioskowania – inference) z centralnej chmury bezpośrednio na urządzenia krańcowe (brzegowe), takie jak czujniki przemysłowe, kamery, smartfony czy maszyny. Przetwarzanie odbywa się jak najbliżej fizycznego źródła danych. Często realizowane jest to w ramach trójwarstwowej architektury: warstwa Edge (urządzenia końcowe) wykonuje zadania wymagające natychmiastowej reakcji; warstwa Fog (lokalne bramy lub serwery) agreguje dane i wykonuje bardziej złożone obliczenia; a warstwa Cloud służy do centralnego trenowania modeli, długoterminowej analityki i globalnego zarządzania. Główne korzyści to minimalna latencja, co jest kluczowe dla aplikacji czasu rzeczywistego (np. autonomiczne pojazdy, robotyka), zwiększona niezawodność (system może działać bez stałego połączenia z internetem), oszczędność pasma sieciowego i poprawa prywatności danych. Edge AI jest technologią napędową dla Przemysłu 4.0, umożliwiając zastosowania takie jak predykcyjne utrzymanie ruchu, wizyjna kontrola jakości w czasie rzeczywistym czy inteligentna automatyzacja (Ali & Dornaika, 2025; Singh & Gill, 2023);
- Architektury hybrydowe - w praktyce, coraz rzadziej stosuje się jeden, czysty paradygmat architektoniczny. Zamiast tego, organizacje budują systemy hybrydowe, które łączą zalety różnych podejść w celu zrównoważenia wydajności, prywatności i skalowalności. Przykładem może być system, w którym wstępne przetwarzanie danych i wnioskowanie w czasie rzeczywistym odbywa się na urządzeniach brzegowych (Edge AI), a następnie zagregowane i zanonimizowane dane są przesyłane do chmury w celu trenowania bardziej złożonych, globalnych modeli. Innym przykładem jest hybryda uczenia maszynowego i sfederowanego, gdzie większość modelu trenowana jest centralnie na

danych publicznych, a następnie dostrajana na urządzeniach klienckich przy użyciu ich prywatnych danych. Badania wskazują, że takie modele hybrydowe konsekwentnie przewyższają wydajnością samodzielne podejścia, choć ich implementacja i zarządzanie nimi pozostają istotnym wyzwaniem.

Poniższa tabela 1 przedstawia syntetyczne porównanie przedstawionych paradygmatów architektonicznych.

Tabela 1. Porównanie paradygmatów architektonicznych AI

Kryterium	Architektura Scentralizowana (Chmurowa)	Uczenie Sfederowane (Federated Learning)	Architektura Brzegowa (Edge AI)
Lokalizacja danych	Scentralizowana (chmura publiczna/prywatna)	Zdecentralizowana (na urządzeniach/w silosach)	Zdecentralizowana (na urządzeniach brzegowych)
Przetwarzanie	Centralne, w potężnych centrach danych	Lokalne trenowanie, centralna agregacja modeli	Lokalne, na urządzeniach o ograniczonej mocy
Prywatność danych	Wysokie ryzyko; wymaga silnych zabezpieczeń i zaufania do dostawcy	Wysoka; surowe dane nie opuszczają lokalizacji	Wysoka; dane przetwarzane lokalnie
Latencja	Wysoka (zależna od sieci)	Umiarkowana do wysokiej (zależna od komunikacji)	Bardzo niska (przetwarzanie w czasie rzeczywistym)
Skalowalność	Bardzo wysoka (elastyczność chmury)	Wysoka (liczba klientów), ale złożona w zarządzaniu	Zależna od liczby i mocy urządzeń brzegowych
Typowe zastosowania	Analiza Big Data, trenowanie dużych modeli językowych, usługi B2C	Sektor medyczny, finanse, współpraca międzyorganizacyjna	Przemysł 4.0, pojazdy autonomiczne, inteligentne miasta, monitoring

Źródło: Opracowanie własne na podstawie: Singh, R., & Gill, S. S. (2023). Edge AI: A survey. *Internet of Things and Cyber-Physical Systems*, 3, 71–92; Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., & Bhagoji, A. N. (2021). *Advances and Open Problems in Federated Learning*. <https://arxiv.org/abs/1912.04977>

Ewolucja od architektur w pełni scentralizowanych w kierunku modeli zdecentralizowanych i brzegowych nie jest jedynie technologiczną modą. Jest to fundamentalna i nieunikniona odpowiedź na ekonomiczne i fizyczne ograniczenia

narzucane przez zjawisko „grawitacji danych”. Koncepcja ta zakłada, że zbiory danych stają się tak duże i kosztowne w przemieszczaniu, że łatwiej i taniej jest przenieść moc obliczeniową do danych, niż dane do mocy obliczeniowej (Sedlak et al., 2023). Architektury takie jak uczenie sfederowane i Edge AI są techniczną realizacją tej zasady. Uczenie sfederowane przenosi proces „trenowania” do danych, podczas gdy Edge AI przenosi proces „wnioskowania” do danych. Dla przedsiębiorstwa, zwłaszcza w dynamicznym otoczeniu SSE, oznacza to, że strategia infrastrukturalna AI musi być projektowana wokół strategii danych, a nie odwrotnie. Decyzja o tym, gdzie dane są generowane i gdzie muszą rezydować (ze względów prawnych, operacyjnych lub kosztowych), bezpośrednio determinuje wybór odpowiedniej architektury. Prawdziwa kompetencja infrastrukturalna polega zatem na zdolności do projektowania i zarządzania hybrydowymi systemami, które elastycznie łączą potęgę centralnej chmury z inteligencją rozproszoną na brzegu sieci, optymalizując przepływ pracy w obliczu nieuniknionej grawitacji danych (Industrial Internet Consortium, 2019).

Analiza usług zarządzania danymi i architektur AI jednoznacznie wskazuje, że sukces wdrożenia sztucznej inteligencji w przedsiębiorstwie zależy od holistycznego, zintegrowanego podejścia. Zaawansowane zarządzanie cyklem życia danych oraz elastyczna, świadomie dobrana do celów biznesowych architektura systemowa, stanowią dwa nierozzerwalne filary gotowości technologicznej. Zaniechanie któregośkolwiek z tych obszarów – czy to poprzez tolerowanie niskiej jakości danych, czy wybór nieodpowiedniej architektury – nieuchronnie prowadzi do niepowodzenia projektów AI, marnotrawstwa zasobów i utraty zaufania do technologii (McGregor & Hostetler, 2023; Singh, 2023).

Dla przedsiębiorstw działających w innowacyjnych i konkurencyjnych ekosystemach, takich jak Specjalne Strefy Ekonomiczne, zdolność do budowania i zarządzania tymi kompetencjami nie jest jedynie opcją, lecz fundamentalnym warunkiem przetrwania i rozwoju. W środowisku, gdzie szybkość innowacji i efektywność operacyjna decydują o pozycji rynkowej, inwestycje w solidny łańd danych oraz nowoczesne, często hybrydowe architektury, łączące skalowalność chmury z responsywnością przetwarzania brzegowego, stają się inwestycjami w długoterminową przewagę konkurencyjną (Ahuja, 2024).

Patrząc w przyszłość, można zidentyfikować wyraźne trendy, które będą kształtować ten obszar. Należy do nich rosnąca automatyzacja procesów zarządzania danymi, w tym powstawanie samonaprawiających się potoków danych, które autonomicznie wykrywają i korygują problemy z jakością (Tewari, 2025).

Sieci, chmura obliczeniowa i zasoby obliczeniowe

Współczesne organizacje funkcjonują w warunkach permanentnej ewolucji technologicznej, określanej mianem Czwartej Rewolucji Przemysłowej lub Przemysłu 4.0, która obejmuje fundamentalne procesy digitalizacji, automatyzacji oraz robotyzacji wszystkich aspektów działalności wytwórczej i zarządczej. Koncepcja ta, sformułowana po raz pierwszy w 2011 roku, opisuje proces integracji zaawansowanych technologii w jednolity system operacyjny, prowadzący do transformacji istniejących produktów i usług w nową, cyfrową formę. Transformacja cyfrowa jest zatem nieustannym procesem integracji świata rzeczywistego i wirtualnego, który staje się kluczowym motorem innowacji i wymusza głębokie zmiany w strukturze przedsiębiorstw. Cyfryzacja często przybiera charakter transformacji wywrotowej, tworząc nowe wartości dla konsumentów oraz redefiniując dynamikę konkurencyjną na rynku. Funkcjonowanie w tak turbulentnym otoczeniu, zdeterminowanym przez innowacyjne technologie, ewoluujące modele konkurencyjności i zmieniające się regulacje, wymaga od organizacji ponadprzeciętnej elastyczności strategicznej i zdolności adaptacyjnych (Ślusarczyk, 2019).

Tradycyjne postrzeganie infrastruktury IT jako zbioru odseparowanych domen – sieci, serwerów i pamięci masowej – staje się anachroniczne w kontekście wymagań stawianych przez sztuczną inteligencję. Analiza dostępnych rozwiązań rynkowych ukazuje głęboko zintegrowany system, w którym platformy chmurowe pełnią rolę kluczowego mechanizmu dostępowego do wyspecjalizowanych zasobów, takich jak procesory graficzne (GPU). Jednocześnie, zaawansowane technologie sieciowe, w szczególności sieć 5G, stają się krwioobiegiem dla nowych topologii przetwarzania danych, takich jak edge computing, które są niezbędne dla zastosowań AI wrażliwych na opóźnienia. Prawdziwa gotowość technologiczna organizacji jest zatem miarą jej zdolności do strategicznego zarządzania tym zintegrowanym systemem, a nie jedynie biegłości w obsłudze poszczególnych jego komponentów. W erze Przemysłu 4.0, gdzie cyfryzacja i AI stają się kluczowymi czynnikami konkurencyjności, zdolność do harmonijnego łączenia tych elementów decyduje o efektywności operacyjnej i potencjale innowacyjnym przedsiębiorstwa (Stelmach, 2025).

Fundamentalna rola, jaką odgrywają procesory graficzne w rozwoju nowoczesnej sztucznej inteligencji, wynika bezpośrednio z ich unikalnej architektury, która jest immanentnie lepiej przystosowana do wymagań obliczeniowych głębokiego uczenia niż architektura tradycyjnych procesorów. Podczas gdy CPU są zoptymalizowane pod kątem minimalizacji opóźnień dla pojedynczych,

sekwencyjnych zadań, posiadając kilka lub kilkadziesiąt potężnych, wszechstronnych rdzeni, GPU zostały zaprojektowane z myślą o maksymalizacji przepustowości poprzez masowe przetwarzanie równoległe. Ich budowa opiera się na tysiącach znacznie prostszych, wyspecjalizowanych rdzeni, które doskonale sprawdzają się w modelu obliczeniowym SIMD/SIMT (Single Instruction, Multiple Data/Threads), gdzie ta sama operacja matematyczna jest wykonywana jednocześnie na ogromnych zbiorach danych. Taki paradygmat działania idealnie odzwierciedla naturę obliczeń w sieciach neuronowych, które w dużej mierze sprowadzają się do operacji na macierzach i tensorach, takich jak mnożenie macierzy (Ashfaq, 2025).

Struktura sprzętowa GPU, zorganizowana w siatkę bloków obliczeniowych, z których każdy zawiera liczne wątki przetwarzające, pozwala na efektywne rozproszenie i równoległe wykonanie milionów niezależnych kalkulacji niezbędnych do trenowania głębokich modeli. To właśnie ta zdolność do jednoczesnego przetwarzania ogromnych ilości danych sprawia, że procesy, które na CPU trwałyby tygodniami lub miesiącami, na GPU mogą być realizowane w ciągu godzin lub dni. Architektura ta, pierwotnie stworzona na potrzeby renderowania grafiki komputerowej, okazała się kluczowa dla rewolucji w dziedzinie AI, umożliwiając praktyczne zastosowanie złożonych modeli, które wcześniej istniały jedynie w sferze teoretycznej. Różnice w filozofii projektowej obu typów procesorów zostały zilustrowane w poniższej tabeli, która syntetyzuje ich kluczowe cechy w kontekście zastosowań AI (Helmick, 2022; Ravikumar et al., 2022) (tab. 2).

Tabela 2. Porównanie architektur obliczeniowych dla zastosowań AI (CPU vs. GPU)

Cecha	Procesor	Procesor graficzny
Filozofia projektowa	Zoptymalizowany pod kątem jak najniższego opóźnienia dla pojedynczych zadań sekwencyjnych.	Zoptymalizowany pod kątem jak najwyższej przepustowości dla tysięcy zadań równoległych.
Liczba rdzeni	Kilka do kilkadziesiątu wysoce zaawansowanych, potężnych rdzeni.	Setki do tysięcy znacznie prostszych, wyspecjalizowanych rdzeni (np. CUDA Cores).
Model przetwarzania	Przetwarzanie szeregowe i wielozadaniowość (każdy rdzeń może pracować nad innym, złożonym zadaniem).	Masowe przetwarzanie równoległe (SIMD/SIMT): wiele rdzeni wykonuje tę samą instrukcję na różnych danych jednocześnie.

Pamięć	Dostęp do dużej puli pamięci RAM o ogólnym przeznaczeniu, zoptymalizowanej pod kątem szybkości dostępu.	Dedykowana, ultraszybka pamięć VRAM o bardzo dużej przepustowości, kluczowa dla „karmienia” tysięcy rdzeni danymi.
Idealne dla...	Systemy operacyjne, bazy danych, zadania wymagające skomplikowanej logiki i podejmowania decyzji.	Trenowanie głębokich sieci neuronowych, symulacje naukowe, renderowanie grafiki, przetwarzanie dużych zbiorów danych.

Źródło: Opracowanie własne na podstawie: Ashfaq, B. (2025). Why GPUs rule AI and graphics: The ultimate guide to parallel computing. <https://medium.com/@ashfaqbs/why-gpus-rule-ai-and-graphics-the-ultimate-guide-to-parallel-computing-f6109ae166de>; Helmick, L. (2022). General Purpose Computing on Graphics Processing Units for Accelerated Deep Learning in Neural Networks. <https://digitalcommons.liberty.edu/cgi/-viewcontent.cgi?-article=2258&context=honors>; Ravikumar, S., Sriraman, K., Saketh, T., Lokesh, M., & Karanam, S. (2022). Effect of neural network structure in accelerating performance and accuracy of a convolutional neural network with GPU/TPU for image analytics. <https://pmc.ncbi.nlm.nih.gov/articles/PMC9044238/>

Wydajność procesorów graficznych w zastosowaniach AI nie jest jedynie funkcją liczby rdzeni, ale wynikiem ciągłej ewolucji wyspecjalizowanych komponentów sprzętowych oraz dojrzałego ekosystemu oprogramowania, który umożliwia pełne wykorzystanie ich potencjału. Kluczową rolę w tym procesie odgrywa platforma obliczeń równoległych NVIDIA CUDA (Compute Unified Device Architecture), która stanowi krytyczny pomost między sprzętem a frameworkami głębokiego uczenia, takimi jak TensorFlow czy PyTorch. CUDA to nie tylko sterownik, ale kompletna platforma programistyczna z bogatym zestawem bibliotek (np. cuDNN do operacji na sieciach neuronowych, cuBLAS do algorytmów liniowej), która abstrahuje złożoność architektury GPU i udostępnia programistom potężne narzędzia do tworzenia wysoce zoptymalizowanych aplikacji (Vijona, 2025).

Równoległe do rozwoju oprogramowania, postępowała ewolucja samego sprzętu. Przełomem było wprowadzenie przez firmę NVIDIA wyspecjalizowanych jednostek obliczeniowych, znanych jako Tensor Cores. Są to dedykowane komponenty w architekturze GPU, zaprojektowane do drastycznego przyspieszania operacji mnożenia macierzy, które stanowią rdzeń obliczeń w głębokim uczeniu. Tensor Cores umożliwiają tzw. obliczenia w precyzji mieszanej, gdzie krytyczne operacje mogą być wykonywane z wysoką precyzją (np. 32-bitową, FP32), podczas gdy większość obliczeń odbywa się z precyzją niższą (np. 16-bitową, FP16, lub 8-bitową, FP8). Zastosowanie niższej precyzji przynosi dwójakie

korzyści: po pierwsze, znacząco redukuje zapotrzebowanie na pamięć, a po drugie, zwiększa przepustowość obliczeniową, co bezpośrednio przekłada się na skrócenie czasu trenowania modeli. Najnowsze generacje GPU wprowadzają jeszcze niższe poziomy precyzji, takie jak 4-bitowa (FP4), co dodatkowo potęguje te korzyści bez istotnej utraty dokładności modelu (Maheswaran, 2025).

Chociaż procesory graficzne zdominowały rynek akceleratorów AI, rozwój technologiczny doprowadził do powstania wyspecjalizowanych architektur, takich jak jednostki przetwarzania tensorowego oraz układy FPGA (Field-Programmable Gate Array), które oferują unikalne korzyści w określonych zastosowaniach. TPU, opracowane przez Google, to układy ASIC (Application-Specific Integrated Circuit) zaprojektowane od podstaw w celu maksymalizacji wydajności operacji tensorowych, stanowiących fundament obliczeń w sieciach neuronowych. W przeciwieństwie do GPU, które zachowują pewien stopień uniwersalności, TPU są zoptymalizowane pod kątem jednego zadania: szybkiego i energooszczędnego przetwarzania macierzy. Dzięki tej specjalizacji, w zadaniach takich jak trenowanie dużych modeli językowych, TPU mogą oferować wyższą wydajność i efektywność energetyczną w porównaniu do GPU (Ravikumar et al., 2022).

Z kolei układy FPGA reprezentują zupełnie inne podejście, oferując maksymalną elastyczność i możliwość rekonfiguracji na poziomie sprzętowym. FPGA to matryce programowalnych bramek logicznych, które mogą być konfigurowane do implementacji dowolnego obwodu cyfrowego, co pozwala na stworzenie akceleratora idealnie dopasowanego do konkretnego algorytmu AI. Ta zdolność do „projektowania sprzętu w locie” sprawia, że FPGA doskonale sprawdzają się w zastosowaniach wymagających ultraniskich opóźnień oraz w środowiskach o ograniczonym budżecie energetycznym, takich jak urządzenia brzegowe. Ich główną wadą jest znacznie wyższy próg wejścia – programowanie FPGA wymaga specjalistycznych umiejętności z zakresu projektowania sprzętu. Wybór między GPU, TPU a FPGA nie jest zatem kwestią wskazania „najlepszego” akceleratora, ale strategiczną decyzją uzależnioną od specyfiki zadania, wymagań dotyczących wydajności, opóźnień, kosztów, zużycia energii oraz dostępnych kompetencji w zespole (Wang & Luo, 2022).

Dla firmy działającej w polskiej Specjalnej Strefie Ekonomicznej strategiczna decyzja nie sprowadza się do prostego wyboru modelu GPU, ale do inwestycji w określony ekosystem technologiczny. Taki wybór ma długofalowe implikacje dla procesów rekrutacyjnych (konieczność pozyskania programistów z doświadczeniem w CUDA), kompatybilności oprogramowania oraz dostępu

do najnowszych innowacji w dziedzinie AI, które są w pierwszej kolejności optymalizowane pod kątem tej dominującej platformy. Sprzęt i oprogramowanie są tu nierozdzielnie związane, tworząc ścieżkę zależności, która w istotny sposób kształtuje strategiczne decyzje IT i determinuje długoterminowe zdolności innowacyjne organizacji. W erze generatywnej AI, gdzie poszczególni dostawcy chmurowi optymalizują swoje platformy pod kątem konkretnych modeli (np. ChatGPT na Azure, Gemini na GCP, Claude na AWS), ryzyko vendor lock-in staje się jeszcze bardziej dotkliwie, ograniczając zdolność firmy do swobodnego wyboru i integracji najlepszych dostępnych technologii (Opara-Martins et al., 2016).

Chmura obliczeniowa w sposób fundamentalny redefiniuje ekonomiczne aspekty wdrażania sztucznej inteligencji, demokratyzując dostęp do zaawansowanych zasobów obliczeniowych. Tradycyjny model, wymagający ogromnych inwestycji kapitałowych (CAPEX) w zakup i utrzymanie własnej, wyspecjalizowanej infrastruktury serwerowej, stanowił barierę nie do pokonania dla wielu organizacji, zwłaszcza dla małych i średnich przedsiębiorstw (MŚP) oraz startupów, które często stanowią trzon ekosystemów innowacji w SSE. Chmura wprowadza model oparty na wydatkach operacyjnych (OPEX), w którym organizacje płacą jedynie za faktycznie wykorzystane zasoby w formule pay-as-you-go. Taka zmiana paradygmatu eliminuje ryzyko finansowe związane z początkowymi inwestycjami i pozwala firmom każdej wielkości na eksperymentowanie, prototypowanie i wdrażanie rozwiązań AI bez konieczności zamrażania znacznego kapitału. Redukcja kosztów staje się w tym kontekście strategicznym czynnikiem, który obniża próg wejścia i stymuluje innowacyjność (Chen, 2025).

Poza korzyściami ekonomicznymi, kluczowa wartość chmury obliczeniowej dla zastosowań AI leży w jej niezrównanej elastyczności i skalowalności. Obciążenia związane z AI charakteryzują się dużą zmiennością zapotrzebowania na moc obliczeniową – faza trenowania modelu może wymagać dostępu do setek lub tysięcy rdzeni GPU przez ograniczony czas, podczas gdy faza wnioskowania (inferencji) na gotowym modelu jest znacznie mniej zasobochłonna. Chmura pozwala na dynamiczne, niemal natychmiastowe skalowanie zasobów w górę i w dół, precyzyjnie dopasowując je do aktualnych potrzeb projektu. Ta „hiperskalowalność” eliminuje problem niedostatecznych lub niewykorzystanych zasobów, który jest typowy dla statycznej infrastruktury on-premise. Co więcej, dostawcy chmurowi oferują bogaty ekosystem gotowych do użycia, prekonfigurowanych środowisk, zintegrowanych frameworków (np. TensorFlow, PyTorch) oraz zaawansowanych usług zarządzanych, które abstrahują znaczną część złożoności związanej

z zarządzaniem infrastrukturą. Dzięki temu zespoły deweloperskie mogą skoncentrować się na tworzeniu modeli i logice biznesowej, a nie na czasochłonnej konfiguracji i utrzymaniu sprzętu, co radykalnie skraca czas wdrożenia innowacji (Kostro, 2021).

Ewolucja chmury obliczeniowej doprowadziła do powstania zintegrowanych platform typu AI-as-a-Service (AIaaS), takich jak Amazon SageMaker, Microsoft Azure Machine Learning czy Google Cloud AI Platform. Platformy te stanowią kompleksowe, zarządzane środowiska, które obejmują cały cykl życia projektu uczenia maszynowego – od przygotowania i etykietowania danych, przez budowę, trenowanie i walidację modeli, aż po ich wdrażanie, monitorowanie i skalowanie. AIaaS demokratyzuje dostęp do zaawansowanych narzędzi, oferując zarówno gotowe, pre-trenowane modele (np. do rozpoznawania obrazu, przetwarzania języka naturalnego), jak i środowiska typu AutoML, które automatyzują proces wyboru i optymalizacji algorytmów, umożliwiając tworzenie modeli nawet osobom bez głębokiej wiedzy eksperckiej. Dla przedsiębiorstw w SSE, platformy te stanowią potężny akcelerator innowacji, pozwalając na szybkie prototypowanie i wdrażanie rozwiązań AI bez konieczności budowania od zera skomplikowanej infrastruktury i zatrudniania wysoko wyspecjalizowanych zespołów (Soumya, 2025; Patel et al., 2024). Wybór konkretnej platformy AIaaS jest decyzją strategiczną, która zależy od istniejącej infrastruktury chmurowej organizacji, preferowanych narzędzi oraz specyfiki projektu. Amazon SageMaker wyróżnia się szeroką integracją z ekosystemem AWS oraz bogatym rynkiem gotowych algorytmów. Google Cloud AI Platform czerpie z zaawansowanych badań Google w dziedzinie AI, oferując ścisłą integrację z frameworkiem TensorFlow oraz dostęp do potężnych akceleratorów TPU. Z kolei Microsoft Azure Machine Learning jest często wybierany przez organizacje silnie osadzone w ekosystemie Microsoftu, oferując zaawansowane narzędzia do odpowiedzialnego AI, które pomagają w zapewnieniu zgodności, transparentności i uczciwości modeli. Niezależnie od wyboru, korzystanie z platform AIaaS wiąże się z ryzykiem wspomnianego wcześniej uzależnienia od dostawcy, co wymaga od przedsiębiorstw świadomego projektowania architektury w sposób umożliwiający przyszłą migrację i interoperacyjność (Chinthapatla, 2023).

Polski rynek chmury obliczeniowej wykazuje znaczną dynamikę wzrostu, a jego wartość w 2024 roku osiągnęła 4,7 mld zł, co stanowi wzrost o 24% w stosunku do roku poprzedniego. Prognozy wskazują, że do 2030 roku wartość ta przekroczy 13 mld zł. Dane pokazują również wysoki poziom adopcji chmury wśród dużych firm (88%), jednak w segmencie MŚP wskaźnik ten jest znacznie

niższy i wynosi 53%. Pomimo tego pozytywnego trendu, poziom wykorzystania zaawansowanych usług chmurowych w Polsce wciąż pozostaje w tyle za europejskimi liderami. Analizy wskazują, że głównymi barierami hamującymi głębszą transformację nie są ograniczenia technologiczne, lecz czynniki organizacyjne i ludzkie. Do najczęściej wymienianych należą: deficyt specjalistycznych kompetencji na rynku pracy, obawy dotyczące bezpieczeństwa danych (często wynikające z braku świadomości i wiedzy), a także niepewność regulacyjna, zwłaszcza w sektorach ściśle regulowanych, które obawiają się naruszenia przepisów dotyczących suwerenności i przechowywania danych (PMR, 2025; Kostro, 2022).

Zestawienie danych rynkowych i barier adopcyjnych ujawnia istnienie w Polsce krytycznego paradoksu: z jednej strony kraj posiada wysoki potencjał naukowy i akademicki w dziedzinie AI (Polska zajmuje piąte miejsce w UE pod względem liczby publikacji naukowych o AI), z drugiej zaś strony charakteryzuje się bardzo niskim poziomem wdrożeń AI w przedsiębiorstwach (zaledwie 3,7% firm deklaruje korzystanie z narzędzi AI). Ta rozbieżność wskazuje na istnienie głębokiej „luki wdrożeniowej”. Dla przedsiębiorstw zlokalizowanych w SSE, które z założenia mają być motorami innowacji, ta luka stanowi centralne wyzwanie strategiczne. Problem nie leży w braku dostępu do technologii, ale w deficycie zdolności organizacyjnych, wizji strategicznej i kapitału ludzkiego niezbędnych do jej efektywnego wykorzystania. Wynika z tego, że polityka wsparcia w ramach SSE powinna ewoluować, przenosząc punkt ciężkości z czysto finansowych instrumentów motywacyjnych na programy budujące konkretne kompetencje: szkolenia z zakresu technologii chmurowych, edukację w obszarze cyberbezpieczeństwa oraz doradztwo w zakresie zgodności regulacyjnej (Cichy et al., 2024; Uchwat, 2025).

Piąta generacja sieci komórkowych (5G) stanowi technologiczną rewolucję, która wykracza daleko poza inkrementalny wzrost prędkości transmisji danych. Jest to fundamentalny element umożliwiający rozwój nowej klasy rozproszonych aplikacji AI działających w czasie rzeczywistym. Definiujące cechy sieci 5G – ultrawysoka niezawodność i ultraniskie opóźnienia (uRLLC), masowa komunikacja między maszynami (mMTC) oraz rozszerzona mobilna łączność szerokopasmowa (eMBB) – tworzą swoisty „system nerwowy” dla gospodarki cyfrowej. Infrastruktura ta jest niezbędna do obsługi Internetu Rzeczy (IoT) – zaawansowanej koncepcji technologicznej umożliwiającej zdalną akwizycję, integrację i przetwarzanie danych z rozproszonych źródeł w celu optymalizacji procesów i monitorowania operacyjnego. Synergia między AI, IoT i 5G ma potencjał do transformacji całych sektorów gospodarki, od logistyki po opiekę zdrowotną,

umożliwiając tworzenie systemów, które w czasie rzeczywistym zbierają, analizują i reagują na dane z otoczenia (Zreikat et al., 2025). Dla wielu krytycznych zastosowań sztucznej inteligencji, takich jak sterowanie robotami przemysłowymi czy systemy wspomagania kierowcy, opóźnienie (latencja) związane z przesyłaniem danych do odległego, scentralizowanego centrum danych w chmurze jest nieakceptowalne. Paradygmat edge computing (przetwarzania brzegowego) stanowi odpowiedź na to wyzwanie. Polega on na przeniesieniu zasobów obliczeniowych i zdolności do wnioskowania (inferencji) modeli AI jak najbliżej źródła generowania danych – na przykład bezpośrednio na hali produkcyjnej, wewnątrz pojazdu, czy na stacji bazowej sieci komórkowej. Taka decentralizacja obliczeń minimalizuje opóźnienia, redukuje obciążenie sieci szkieletowej oraz zwiększa prywatność i bezpieczeństwo, ponieważ wrażliwe dane mogą być przetwarzane lokalnie, bez konieczności ich transferu na zewnątrz. Edge computing jest zatem kluczowym uzupełnieniem architektury chmurowej, niezbędnym do pełnego wykorzystania potencjału aplikacji AI działających w czasie rzeczywistym, które są akcelerowane przez sieć 5G (Thota, 2024).

Przyszłość infrastruktury AI w przedsiębiorstwach nie będzie opierać się na binarnym wyborze między przetwarzaniem brzegowym a chmurowym, lecz na hybrydowym, dynamicznym kontinuum, w którym oba te modele współistnieją i wzajemnie się uzupełniają. W tej architekturze obciążenia obliczeniowe są inteligentnie orkiestrowane w całej infrastrukturze. Kluczowym elementem jest tu warstwa orkiestracji, która dynamicznie decyduje, gdzie i jak przetwarzać dane oraz wykonywać zadania AI. Decyzje te opierają się na szeregu czynników, takich jak wrażliwość na opóźnienia, złożoność modelu, dostępność przepustowości sieciowej oraz wymogi dotyczące prywatności i suwerenności danych. Zasobochłonne procesy trenowania modeli odbywają się w potężnych, scentralizowanych centrach danych w chmurze, podczas gdy szybkie i zwinne wnioskowanie w czasie rzeczywistym ma miejsce na brzegu sieci. Taka synergia pozwala na optymalne wykorzystanie zasobów i maksymalizację wydajności (John, 2025).

Architektoniczna ewolucja w kierunku kontinuum edge-cloud jest strategiczną odpowiedzią na rosnące wyzwania związane z zarządzaniem danymi, prywatnością i odpornością systemów. Przetwarzanie danych lokalnie, na brzegu sieci, pozwala przedsiębiorstwom na wdrażanie zaawansowanych rozwiązań AI przy jednoczesnym zachowaniu zgodności z rygorystycznymi regulacjami, takimi jak RODO/GDPR. Zwiększa to również odporność operacyjną, gdyż systemy brzegowe mogą kontynuować działanie autonomicznie nawet w przypadku utraty łączności z chmurą. Dla przedsiębiorstw z SSE, zwłaszcza tych działających

w sektorach regulowanych (np. technologii medycznych, finansów), przyjęcie architektury edge-cloud może być nie tylko optymalizacją, ale warunkiem koniecznym do wdrożenia innowacji opartych na AI. Zdolność do projektowania systemów, które są nie tylko szybsze, ale także bardziej bezpieczne, odporne i zgodne z prawem, staje się kluczowym źródłem przewagi konkurencyjnej (Belcastro et al., 2025; Sotonye et al., 2025).

Patrząc w przyszłość, koncepcje związane z siecią 6G zakładają jeszcze głębszą integrację AI z samą tkanką sieciową. Sieć 6G jest projektowana jako „AI-native”, co oznacza, że sztuczna inteligencja i uczenie maszynowe będą wbudowane w jej architekturę jako podstawowa zdolność, umożliwiając sieci autonomiczne zarządzanie zasobami, dynamiczną optymalizację i świadczenie rozproszonych usług AI jako integralnej części swojej funkcjonalności. W takim paradygmacie sieć nie jest już tylko pasywnym medium transmisyjnym, ale staje się aktywnym, inteligentnym systemem obliczeniowym. To przejście od scentralizowanej do sfederowanej, rozproszonej inteligencji ma fundamentalne znaczenie nie tylko technologiczne, ale i biznesowe, otwierając drogę do nowych, zdecentralizowanych aplikacji i modeli biznesowych, które dziś są trudne do wyobrażenia (Frąckiewicz, 2025).

Rzeczywista gotowość technologiczna przedsiębiorstw do wdrożenia i efektywnego wykorzystania sztucznej inteligencji jest uwarunkowana strategiczną integracją trzech kluczowych komponentów infrastrukturalnych: wyspecjalizowanych zasobów obliczeniowych, skalowalnych platform chmurowych oraz zaawansowanych sieci telekomunikacyjnych. Te trzy filary nie funkcjonują w izolacji, lecz tworzą złożony, synergiczny ekosystem, którego opanowanie jest fundamentalnym wyzwaniem i jednocześnie największą szansą dla nowoczesnych organizacji. Dla przedsiębiorstw działających w ramach polskich Specjalnych Stref Ekonomicznych, które z definicji mają pełnić rolę krajowych centrów innowacji i konkurencyjności, mistrzowskie zarządzanie tym ekosystemem staje się nie tyle celem technicznym, co głównym determinantem ich zdolności do konkurowania na globalnym rynku. Ambicje Polski, by stać się jednym z kluczowych ośrodków AI w Europie, wymagają stworzenia dynamicznego i wspierającego ekosystemu, w którym firmy mogą rozwijać i wdrażać przełomowe technologie (Stathopoulou et al., 2025; Kokkonen et al., 2022).

Kluczowym wnioskiem płynącym jest identyfikacja „luki wdrożeniowej” jako centralnej bariery dla rozwoju AI w Polsce. Zjawisko to można zinterpretować w ramach koncepcji dojrzałości cyfrowej, która określa zaawansowanie procesu transformacji cyfrowej, obejmując zarówno dotychczasowe osiągnięcia

technologiczne, jak i rozwój kompetencji zarządczych. Pomimo silnego zaplecza naukowo-badawczego, polskie przedsiębiorstwa wciąż w ograniczonym stopniu adoptują technologie chmurowe i rozwiązania oparte na AI, co świadczy o niskiej dojrzałości cyfrowej. Jest ona hamowana przez deficyt kompetencji, obawy o bezpieczeństwo i niepewność regulacyjną. Przewycięzenie tego wyzwania wymaga holistycznego podejścia, które oprócz inwestycji w infrastrukturę (zasoby cyfrowe) obejmuje również rozwój zdolności organizacji do transformacji, wyrażającej się poprzez skuteczne formułowanie wizji strategicznej, budowanie adekwatnej kultury organizacyjnej, wzmacnianie przywództwa oraz rozwijanie kompetencji w zakresie zarządzania zmianą i innowacyjnością. Polityka rozwoju AI w Polsce, oparta na filarach kapitału ludzkiego, innowacji, inwestycji i wdrożeń, musi znaleźć swoje praktyczne odzwierciedlenie w działaniach podejmowanych w ramach SSE, aby mogły one w pełni zrealizować swoją misję bycia awangardą polskiej gospodarki cyfrowej.

Bezpieczeństwo danych jako fundament infrastruktury AI

Dane jako kluczowy zasób napędzający algorytmy i modele uczenia maszynowego, determinują jakość, precyzję i wiarygodność wyników generowanych przez te systemy. Jak trafnie ujmują to Russell i Norvig, sztuczna inteligencja jest w swojej istocie procedurą rozwiązywania problemów, która generuje wyniki poprzez analizę danych wejściowych; dane są paliwem dla całego procesu (Russell & Norvig, 2022). Wszelkie naruszenia ich integralności, poufności czy dostępności prowadzą nie tylko do bezpośredniego ryzyka operacyjnego, ale podważają całą wartość biznesową i strategiczną inwestycji w AI.

Paradoksu bezpieczeństwa AI polega na rosnącej dysproporcji między wysoką świadomością zagrożeń a nieadekwatnym poziomem implementacji zabezpieczeń w organizacjach. Dane przedstawione w raporcie World Economic Forum na rok 2025 są alarmujące: podczas gdy 66% organizacji spodziewa się, że sztuczna inteligencja będzie miała najbardziej znaczący wpływ na krajobraz cyberbezpieczeństwa w nadchodzącym roku, zaledwie 37% z nich deklaruje posiadanie wdrożonych, odpowiednich procesów ochronnych. Ta luka nie jest jedynie statystyczną anomalią, lecz stanowi strategiczne wyzwanie, szczególnie dla przedsiębiorstw funkcjonujących w ramach Specjalnych Stref Ekonomicznych (SSE), których długoterminowa konkurencyjność coraz silniej zależy od zdolności do bezpiecznej i efektywnej adopcji technologicznych innowacji (Jurgens & Dal Cin, 2025).

Wspomniany paradoks staje się motorem napędowym systemowego ryzyka, które manifestuje się poprzez zjawisko „cybernierówności”. Wdrożenie zaawansowanych rozwiązań AI wymaga nie tylko znaczących inwestycji w moc obliczeniową i oprogramowanie, ale również w wysoce specjalistyczną infrastrukturę bezpieczeństwa, której koszty mogą być barierą nie do pokonania dla mniejszych podmiotów (Jurgens & Dal Cin, 2025). W rezultacie duże, zasobne korporacje są w stanie absorbować te koszty, implementując zarówno innowacyjne systemy AI, jak i solidne mechanizmy ich ochrony. Jednocześnie małe i średnie przedsiębiorstwa (MŚP), często działające pod presją konkurencyjną, wdrażają bardziej dostępne, np. chmurowe, narzędzia AI, nie posiadając jednak wystarczających zasobów lub kompetencji do ich właściwego zabezpieczenia. Taka sytuacja prowadzi do powstania systemowej podatności, ponieważ MŚP są kluczowymi ogniwami w złożonych łańcuchach dostaw – problem ten jest postrzegany jako największe wyzwanie dla cyberodporności przez 54% dużych organizacji. Niezabezpieczone MŚP staje się zatem potencjalnym wektorem ataku na cały ekosystem gospodarczy, czyniąc z „paradoksu bezpieczeństwa AI” nie tylko problem pojedynczej firmy, ale zagrożenie dla stabilności całej strefy ekonomicznej (Jurgens & Dal Cin, 2025).

Dynamiczny rozwój sztucznej inteligencji doprowadził do fundamentalnej zmiany w krajobrazie cyberbezpieczeństwa, nadając jej podwójną, antagonistyczną rolę. Z jednej strony, AI stała się potężnym narzędziem w rękach cyberprzestępców, umożliwiając ataki o niespotykanej dotąd skali i wyrafinowaniu. Z drugiej strony, te same technologie stanowią trzon nowoczesnych, proaktywnych strategii obronnych, pozwalając organizacjom wyprzedzać działania adversarzy. Zrozumienie tej dwoistej natury jest kluczowe dla budowy odpornej infrastruktury (Jeffrey et al., 2025).

Cyberprzestępcy aktywnie wykorzystują modele AI, w szczególności generatywnej, do przełamywania tradycyjnych barier obronnych. Zamiast statycznych, łatwych do zidentyfikowania metod, stosują oni dynamiczne, uczące się algorytmy, które adaptują się do środków zaradczych w czasie rzeczywistym. Główne obszary wykorzystania AI w działaniach ofensywnych obejmują zaawansowaną inżynierię społeczną, automatyzację ataków oraz wykorzystanie technologii deepfake do celów dezinformacyjnych i szpiegowskich. Globalne koszty cyberprzestępczości, napędzane m.in. przez ataki wspomagane przez AI, mają osiągnąć astronomiczną kwotę 10,5 biliona dolarów rocznie do 2025 roku, co obrazuje skalę problemu (Łochowska, 2025). Generatywna sztuczna inteligencja zrewolucjonizowała inżynierię społeczną, umożliwiając tworzenie masowych, a jednocześnie wysoce

spersonalizowanych i lingwistycznie poprawnych kampanii phishingowych. Modele językowe potrafią automatycznie generować przekonujące wiadomości e-mail, które naśladują styl komunikacji konkretnych osób lub instytucji, co czyni je niezwykle trudnymi do odróżnienia od legalnej korespondencji. Podobnie, techniki syntezy mowy (tzw. vishing) pozwalają na tworzenie fałszywych wiadomości głosowych, np. od przełożonego proszącego o pilny przelew, co stanowi poważne zagrożenie dla bezpieczeństwa finansowego przedsiębiorstw (Łochowska, 2025).

Algorytmy uczenia maszynowego są również wykorzystywane do pełnej automatyzacji różnych faz cyberataku. Systemy oparte na AI potrafią samodzielnie skanować sieci korporacyjne w poszukiwaniu luk w zabezpieczeniach, testować różne wektory eksploatacji i adaptacyjnie omijać systemy detekcji intruzów. Duże modele językowe (LLM) są już stosowane do analizy kodu źródłowego w celu identyfikacji podatności oraz do automatycznego generowania złośliwego oprogramowania i skryptów wykorzystujących te luki (Panfil, 2025; Xu et al., 2025). Szczególnie niebezpieczne jest wykorzystanie technologii deepfake, która pozwala na tworzenie realistycznych, fałszywych materiałów wideo i audio. W kontekście biznesowym może być ona użyta do szpiegostwa korporacyjnego, na przykład poprzez podszywanie się pod członków zarządu podczas videokonferencji w celu wyłudzenia poufnych informacji, a także do manipulacji reputacją firmy lub destabilizacji rynków finansowych (Panfil, 2025).

Równolegle do ofensywnego wykorzystania AI, technologie te stają się niezastąpionym elementem nowoczesnych strategii cyberobrony. Tradycyjne, reaktywne centra operacji bezpieczeństwa (SOC), które polegają na manualnej analizie alertów, ewoluują w kierunku proaktywnych, predykcyjnych jednostek (AI-SOC), zdolnych do identyfikacji i neutralizacji zagrożeń, zanim dojdzie do eskalacji incydentu. Wykorzystanie AI w obronie pozwala na przetwarzanie i korelowanie ogromnych wolumenów danych w czasie rzeczywistym, co jest niemożliwe do osiągnięcia dla ludzkich analityków. Przewiduje się, że rynek cyberbezpieczeństwa oparty na generatywnej AI osiągnie wartość 35,5 miliarda dolarów do 2031 roku, co świadczy o rosnącym znaczeniu tych technologii w strategiach obronnych (Research and Markets, 2025).

Jednym z kluczowych zastosowań AI w cyberbezpieczeństwie jest inteligentne wykrywanie anomalii. Systemy uczenia maszynowego, w szczególności głębokie sieci neuronowe, są trenowane na historycznych danych telemetrycznych (logach systemowych, metrykach wydajności, ruchu sieciowym), aby nauczyć się wzorców normalnego zachowania infrastruktury IT. Każde subtelne odchylenie od tej normy jest natychmiast flagowane jako potencjalny incydent

bezpieczeństwa. Takie podejście pozwala na wykrywanie zaawansowanych, nieznanych wcześniej zagrożeń oraz ataków typu Advanced Persistent Threat (APT), które umykają tradycyjnym systemom opartym na sygnaturach (Azizi et al., 2024; Jeffrey et al., 2023). Sztuczna inteligencja odgrywa również centralną rolę w platformach typu SOAR (Security Orchestration, Automation, and Response). Po wykryciu anomalii, systemy AI mogą automatycznie uruchomić predefiniowane scenariusze reagowania, takie jak izolacja zainfekowanego hosta od sieci, zablokowanie złośliwego adresu IP na firewallu czy unieważnienie skompromitowanych poświadczeń. Automatyzacja tych procesów drastycznie skraca czas reakcji (Mean Time to Recovery - MTTR) i minimalizuje potencjalne szkody, jednocześnie redukując zjawisko „zmęczenia alertami” u analityków bezpieczeństwa, pozwalając im skupić się na bardziej złożonych zadaniach dochodzeniowych.

Zaawansowane modele AI są także wykorzystywane do predykcyjnej analizy zagrożeń. Poprzez analizę globalnych trendów w cyberprzestępczości, dyskusji na forach w darknecie oraz raportów o nowych podatnościach, systemy AI mogą przewidywać, jakie wektory ataków i techniki będą najprawdopodobniej używane w przyszłości przeciwko danej organizacji lub branży. Taka wiedza pozwala na proaktywne wzmacnianie systemów obronnych, wdrażanie odpowiednich poprawek bezpieczeństwa i dostosowywanie polityk, zanim dojdzie do faktycznego ataku (Alevizos & Dekker, 2024). Infrastruktura sztucznej inteligencji, w odróżnieniu od tradycyjnych systemów IT, wprowadza nowy, unikalny zestaw podatności, które wykraczają poza klasyczne zagrożenia sieciowe czy aplikacyjne. Jej powierzchnia ataku obejmuje nie tylko serwery, sieci i bazy danych, ale przede wszystkim kluczowe komponenty cyklu życia modelu AI: dane treningowe, sam algorytm oraz potoki przetwarzania, które je łączą. Zrozumienie tych specyficznych wektorów jest niezbędne do zaprojektowania skutecznej strategii obronnej, ponieważ tradycyjne środki bezpieczeństwa mogą okazać się wobec nich niewystarczające (Zhao et al., 2025).

Jednym z najbardziej podstępnych zagrożeń jest zatrutowanie danych. Atak ten polega na celowym wprowadzeniu subtelnie zmanipulowanych lub szkodliwych danych do zbioru treningowego modelu. Celem może być sabotaż jego działania, na przykład poprzez nauczenie go błędnej klasyfikacji określonych obiektów, co w systemach autonomicznych pojazdów mogłoby prowadzić do katastrofalnych skutków (Panfil, 2025). Innym celem może być wprowadzenie ukrytej „tylnej furtki”, która aktywuje się tylko po otrzymaniu specyficznego, spreparowanego sygnału wejściowego. Co więcej, zatrutowanie danych może prowadzić do generowania przez model stronniczych i dyskryminujących wyników,

co naraża organizację na poważne ryzyka prawne i reputacyjne (Mengara et al., 2023).

Kolejnym krytycznym wektorem są ataki adversarialne, wymierzone we wdrożone już, działające modele. Polegają one na tworzeniu specjalnie spreparowanych danych wejściowych, które dla ludzkiego oka wydają się normalne i niezmienione, ale są zaprojektowane tak, aby wprowadzić model w błąd. Przykładowo, minimalna, niewidoczna dla człowieka modyfikacja obrazu znaku „stop” może spowodować, że system rozpoznawania obrazu w autonomicznym pojeździe sklasyfikuje go jako znak ograniczenia prędkości. Ataki te są uważane za jedno z największych i najtrudniejszych do wykrycia zagrożeń, ponieważ wykorzystują fundamentalne słabości matematyczne modeli głębokiego uczenia. Ich skuteczność podważa zaufanie do wszelkich autonomicznych systemów decyzyjnych, od systemów wykrywania fraudów finansowych po diagnostykę medyczną (Zhao et al., 2025).

Infrastruktura AI jest również narażona na kradzież własności intelektualnej i naruszenie prywatności poprzez specyficzne techniki ataków. Ataki ekstrakcji modelu polegają na tym, że przeciwnik, nie mając dostępu do samego modelu, jest w stanie odtworzyć jego architekturę i parametry poprzez wielokrotne wysyłanie zapytań do jego publicznego API i analizowanie zwracanych odpowiedzi. Skuteczna ekstrakcja jest równoznaczna z kradzieżą cennego, często kosztownego w opracowaniu, zasobu intelektualnego firmy (Zhao et al., 2025). Z kolei ataki inwersji modelu stanowią poważne zagrożenie dla prywatności. Ich celem jest odtworzenie wrażliwych danych, które zostały użyte do treningu modelu, na podstawie jego wyników. Badania wykazały, że z modelu rozpoznawania twarzy można zrekonstruować wizerunki osób ze zbioru treningowego, co stanowi rażące naruszenie ochrony danych osobowych (Zhao et al., 2025). Wreszcie, rosnące uzależnienie od komponentów firm trzecich rodzi poważne zagrożenia w łańcuchu dostaw AI. Współczesne systemy AI rzadko są budowane od zera. Przedsiębiorstwa masowo wykorzystują pre-trenowane modele z publicznych repozytoriów (np. Hugging Face), biblioteki open-source (np. TensorFlow, PyTorch) oraz usługi chmurowe (Model-as-a-Service). Każdy z tych elementów może stać się wektorem ataku. Skompromitowany model pobrany z niezaufanego źródła, zainfekowana biblioteka czy podatność w API dostawcy usług chmurowych mogą umożliwić atakującemu wprowadzenie złośliwego kodu, kradzież danych lub przejęcie kontroli nad całą infrastrukturą AI przedsiębiorstwa. Złożoność i brak transparentności w tych łańcuchach dostaw czynią je jednym z najpoważniejszych wyzwań dla cyberodporności (Ding et al., 2025).

Tabela 3 syntetyzuje omówione wektory ataków specyficznych dla systemów AI, wskazując na ich potencjalny wpływ na działalność przedsiębiorstwa oraz rekomendowane techniczne i organizacyjne środki zaradcze. Stanowi ona skondensowane narzędzie analityczne, które mapuje konkretne ryzyka na adekwatne mechanizmy obronne, tworząc podstawy dla budowy kompleksowej strategii bezpieczeństwa AI.

Skuteczna obrona przed zróżnicowanymi i zaawansowanymi zagrożeniami wymaga wdrożenia wielowarstwowej strategii technologicznej, która zabezpiecza każdy element cyklu życia systemu AI. Poniżej przedstawiono kluczowe filary, na których powinna opierać się odporna infrastruktura AI: integracja bezpieczeństwa z procesem rozwoju (AI SecDevOps) oraz implementacja technologii zwiększających prywatność (Privacy-Enhancing Technologies). Ewolucja, jaka dokonała się w infrastrukturze AI w latach 2019-2025, jest niczym innym jak jej industrializacją. Przeszliśmy od etapu, w którym badacze manualnie kopiowali zbiory danych między serwerami, do w pełni zautomatyzowanych, skonteneryzowanych i wersjonowanych potoków MLOps, obsługiwanych przez platformy takie jak Kubeflow czy TensorFlow Extended (TFX) Google'a. Ta transformacja, odzwierciedlająca wcześniejszą ewolucję w świecie tworzenia oprogramowania od chaotycznego kodowania do ustrukturyzowanych praktyk DevOps, niesie ze sobą fundamentalną implikację dla bezpieczeństwa. Tak jak w DevOps próba „doklejenia” zabezpieczeń na samym końcu procesu okazała się katastrofalnie nieskuteczna, prowadząc do powstania metodyki DevSecOps, tak samo w świecie AI bezpieczeństwo nie może być jedynie audytem gotowego modelu. Musi ono stać się zautomatyzowaną, integralną częścią całego potoku rozwojowego.

Podejście AI SecDevOps zakłada integrację praktyk i narzędzi bezpieczeństwa na każdym etapie cyklu życia AI. Już na etapie projektowania i kodowania, kluczowe staje się wykorzystanie narzędzi do statycznej analizy kodu (Static Application Security Testing - SAST). Rosnący rynek tych rozwiązań, zdolnych do skanowania nie tylko kodu pisanego przez deweloperów, ale również tego generowanego przez modele AI (np. Copilot), a także do analizy zależności w bibliotekach uczenia maszynowego pod kątem znanych podatności, potwierdza ten trend. Wdrożenie SAST w potokach CI/CD (Continuous Integration/Continuous Deployment) pozwala na wczesne wykrywanie i eliminowanie luk, zanim trafią one do środowiska produkcyjnego (Fu et al., 2024).

Tabela 3. Macierz zagrożeń i środków zaradczych w ekosystemach AI

Typ zagrożenia (threat type)	Opis i wpływ na działalność	Techniczne i organizacyjne środki zaradcze
Zatrutowanie danych (Data Poisoning)	Celowe uszkodzenie zbioru treningowego w celu sabotażu wydajności modelu lub wprowadzenia ukrytych słabości. Wpływ: Błędne decyzje biznesowe, ryzyka reputacyjne, dyskryminacja.	Weryfikacja pochodzenia danych (provenance), wykrywanie anomalii w danych wejściowych, ograniczony dostęp do potoków danych, regularne audyty zbiorów danych.
Ataki adversarialne (Adversarial Attacks)	Manipulacja danymi wejściowymi w celu oszukania wdrożonego modelu. Wpływ: Obejście systemów bezpieczeństwa (np. detekcji fraudów), błędna klasyfikacja, utrata zaufania do systemu.	Trening adversarialny, techniki „oczyszczania” danych wejściowych (input sanitization), monitorowanie predykcji modelu pod kątem nietypowych zachowań.
Ekstrakcja modelu (Model Extraction)	Odtworzenie architektury i parametrów modelu poprzez analizę jego odpowiedzi na zapytania. Wpływ: Kradzież własności intelektualnej, utrata przewagi konkurencyjnej.	Ograniczanie liczby zapytań do API, monitorowanie wzorców zapytań, wdrażanie modeli w zaufanych środowiskach wykonawczych, techniki „watermarkingu” modeli.
Ataki na łańcuch dostaw AI	Wykorzystanie skompromitowanych modeli, bibliotek lub API od zewnętrznych dostawców. Wpływ: Wprowadzenie złośliwego kodu do infrastruktury, kradzież danych, przejęcie kontroli nad systemami.	Weryfikacja pochodzenia i integralności modeli, audyty bezpieczeństwa dostawców, izolacja środowisk wykonawczych, SAST.

Źródło: Opracowanie własne na podstawie: Zhao, P., Zhu, W., Jiao, P., Gao, D., & Wu, O. (2025). Data poisoning in deep learning: A survey. <https://arxiv.org/pdf/2503.22759>; Panfil, K. (2025). Zagrożenia sztucznej inteligencji – przykłady i wyjaśnienia: Co musisz wiedzieć w 2025 roku. <https://ttms.com/pl/wyjasnienie-zagrozen-zwiazanych-ze-sztuczna-inteligencja-co-musisz-wiedziec-w-2025-roku/>; Ding, Z., Fu, Q., Ding, J., Deng, G., Liu, Y., & Li, Y. (2025). A rusty link in the AI supply chain: Detecting evil configurations in model repositories. <https://arxiv.org/html/2505.01067v1>

Fundamentalnym elementem dojrzałej infrastruktury MLOps, a co za tym idzie AI SecDevOps, jest wdrożenie kontroli wersji nie tylko dla kodu, ale również dla danych i modeli. Narzędzia takie jak DVC (Data Version Control) pozwalają na traktowanie zbiorów danych z taką samą precyzją jak kod źródłowy, umożliwiając śledzenie zmian, eksperymentów i pochodzenia danych. Taka architektura zapewnia pełną odtwarzalność i audytowalność całego procesu treningowego. W przypadku wykrycia problemu bezpieczeństwa, na przykład ataku

typu data poisoning, organizacja jest w stanie precyzyjnie zidentyfikować skażony zbiór danych i szybko przywrócić model do bezpiecznej, wcześniejszej wersji, minimalizując szkody (Stone et al., 2025; Pacheco et al., 2024).

PETs to klasa zaawansowanych technologii, które w sposób fundamentalny zmieniają podejście do przetwarzania danych. Ich celem jest umożliwienie organizacjom wydobywania cennych informacji i budowania dokładnych modeli AI przy jednoczesnej matematycznej minimalizacji ryzyka naruszenia prywatności osób, których dane dotyczą. W dobie regulacji takich jak RODO, PETs stają się nie tylko narzędziem obronnym, ale strategicznym czynnikiem umożliwiającym innowacje w zgodzie z prawem i oczekiwaniami społecznymi (Boteju et al., 2023; Calvi et al., 2024).

Uczenie sfederowane to paradygmat treningu modeli uczenia maszynowego, który odwraca tradycyjny przepływ danych. Zamiast centralizować ogromne zbiory danych na jednym serwerze w celu trenowania modelu, proces uczenia jest dystrybuowany do miejsc, gdzie dane naturalnie rezydują (np. na urządzeniach mobilnych, w oddziałach szpitalnych, w różnych fabrykach). Każdy lokalny węzeł trenuje kopię modelu na swoich prywatnych danych, a następnie do centralnego serwera przesyłane są jedynie zaktualizowane parametry (wagi) modelu, a nie surowe, wrażliwe dane. Centralny serwer agreguje te parametry w celu stworzenia ulepszanego modelu globalnego, który jest następnie rozsyłany z powrotem do węzłów. Takie podejście fundamentalnie minimalizuje ryzyko masowych wycieków danych i jest kluczowe w sektorach operujących na danych wrażliwych, jak opieka zdrowotna czy finanse (Boteju et al., 2023; Calvi et al., 2024).

Prywatność różnicowa (Differential Privacy) dostarcza formalnej, matematycznej gwarancji prywatności. Jej podstawowa idea polega na dodawaniu precyzyjnie skalibrowanego „szumu” statystycznego do wyników zapytań kierowanych do bazy danych lub w kontekście AI, do parametrów (gradientów) modelu podczas procesu treningu. Ilość dodanego szumu jest na tyle duża, aby uniemożliwić atakującemu stwierdzenie z dużą pewnością, czy dane konkretnej osoby były częścią zbioru treningowego, ale jednocześnie na tyle mała, aby nie zniszczyć użyteczności statystycznej wyników. Prywatność różnicowa jest obecnie uważana za złoty standard w anonimizacji i stanowi jedno z najskuteczniejszych zabezpieczeń przed atakami rekonstrukcji danych, takimi jak inwersja modelu (Shojaei et al., 2025).

Szyfrowanie homomorficzne jest jedną z najbardziej przełomowych technik kryptograficznych, często określaną mianem „świętego Graala” bezpiecznego przetwarzania danych. Umożliwia ono wykonywanie dowolnych obliczeń matematycznych (np. sumowania, mnożenia) bezpośrednio na danych zaszyfrowanych,

bez konieczności ich wcześniejszego odszyfrowywania. W kontekście AI oznacza to potencjalną możliwość trenowania modeli uczenia maszynowego na danych, które pozostają w pełni zaszyfrowane przez cały czas, nawet w niezaufanym środowisku, takim jak publiczna chmura obliczeniowa. Chociaż technologia ta wciąż boryka się z wyzwaniami związanymi z ogromnym narzutem obliczeniowym, jej rozwój postępuje dynamicznie i stanowi ona obietnicę najwyższego poziomu bezpieczeństwa danych w przyszłych systemach AI. Najbardziej odporne konfiguracje zabezpieczeń łączą warstwowo prywatność różnicową, uczenie sfederowane i częściowe szyfrowanie homomorficzne, co pozwala blokować obecne ataki przy zachowaniu wysokiej dokładności modelu (Shojaei et al., 2025).

Wdrożenie zaawansowanych technologii obronnych jest warunkiem koniecznym, ale niewystarczającym do zapewnienia kompleksowego bezpieczeństwa danych w infrastrukturze AI. Równie istotne, a często trudniejsze w implementacji, są solidne ramy organizacyjne, prawne i etyczne, określane mianem ładu korporacyjnego w zakresie AI (AI Governance). Bez spójnej strategii zarządzania, nawet najlepsze narzędzia technologiczne pozostaną nieskuteczne. Każdy system sztucznej inteligencji, który przetwarza dane osobowe obywateli Unii Europejskiej, podlega rygorystycznym wymogom Ogólnego Rozporządzenia o Ochronie Danych (RODO). Szczególne znaczenie w kontekście AI ma Artykuł 22 RODO, który reguluje zautomatyzowane podejmowanie decyzji w indywidualnych przypadkach, w tym profilowanie. Zgodnie z tym przepisem, osoba, której dane dotyczą, ma prawo do tego, by nie podlegać decyzji, która opiera się wyłącznie na zautomatyzowanym przetwarzaniu i wywołuje wobec niej skutki prawne lub w podobny sposób istotnie na nią wpływa. Co kluczowe, administrator danych jest zobowiązany do przekazania tej osobie „istotnych informacji o logice, a także o znaczeniu i przewidywanych konsekwencjach takiego przetwarzania”, co bezpośrednio implikuje konieczność zapewnienia wyjaśnialności modeli (Groszkowski, 2023). Dodatkowo, horyzont regulacyjny jest kształtowany przez nadchodzący Akt o Sztucznej Inteligencji (AI Act), który wprowadza podejście oparte na ryzyku. Systemy AI sklasyfikowane jako „wysokiego ryzyka” (np. te stosowane w rekrutacji, ocenie zdolności kredytowej czy infrastrukturze krytycznej) będą podlegać surowym wymogom, w tym konieczności przeprowadzania oceny skutków dla ochrony danych (DPIA) zgodnie z Art. 35 RODO, zapewnienia wysokiej jakości danych treningowych oraz utrzymywania ludzkiego nadzoru. Potencjalny zbieg sankcji finansowych z AI Act i RODO sprawia, że zgodność z prawem staje się kluczowym elementem zarządzania ryzykiem (Grupa Robocza ds. AI przy KPRM, 2023).

Jednym z największych wyzwań dla przedsiębiorstw działających globalnie jest rosnąca fragmentacja regulacyjna. Różnice w przepisach dotyczących ochrony danych i AI między jurysdykcjami (np. UE, USA, Chiny) zmuszają organizacje do budowania elastycznych, ale przez to bardziej skomplikowanych i kosztownych systemów zgodności, co z kolei może poszerzać powierzchnię ataku i zwiększać ryzyko techniczne. Ponad 76% dyrektorów ds. bezpieczeństwa informacji (CISO) wskazuje fragmentację regulacji jako istotną barierę w efektywnym zarządzaniu ryzykiem (Jurgens & Dal Cin, 2025; Themann, 2025). Wiele zaawansowanych modeli AI, zwłaszcza sieci głębokiego uczenia, działa na zasadzie „czarnej skrzynki”, co oznacza, że ich wewnętrzna logika decyzyjna jest nieprzejrzysta i trudna do zinterpretowania nawet dla ich twórców. Ta nieprzejrzystość stanowi fundamentalną barierę dla adopcji technologii AI, ponieważ podważa zaufanie użytkowników i utrudnia pociągnięcie do odpowiedzialności w przypadku błędnego działania systemu (Astrup, 2020; Azizi et al., 2025).

W tym kontekście, wyjaśnialna AI przestaje być jedynie technicznym narzędziem do debugowania modeli, a staje się kluczowym instrumentem z obszaru zarządzania, ryzyka i zgodności. Zdolność do wyjaśnienia, dlaczego model podjął określoną decyzję (np. dlaczego odrzucił wniosek kredytowy lub zidentyfikował transakcję jako oszustwo), jest nie tylko wymogiem prawnym wynikającym z RODO, ale przede wszystkim warunkiem koniecznym do budowania i utrzymania zaufania klientów oraz partnerów biznesowych. Badania pokazują, że nawet jeśli systemy AI są technicznie skuteczniejsze w wykrywaniu anomalii, ich adopcja jest ograniczona, gdy użytkownicy nie rozumieją i nie ufają ich logice działania. Transparentność i wyjaśnialność są zatem krytycznymi determinantami zaufania, które bezpośrednio przekładają się na komercyjną żywotność rozwiązań AI (Paliwal et al., 2025).

Skuteczne zarządzanie ryzykiem związanym z AI wymaga ustanowienia w organizacji formalnych, wewnętrznych ram zarządczych. Powinny one obejmować powołanie interdyscyplinarnych zespołów (łączyących ekspertów z IT, prawa, biznesu i etyki), które będą odpowiedzialne za nadzór nad odpowiedzialnym wdrażaniem i wykorzystaniem technologii AI. Do kluczowych zadań takich zespołów należy przeprowadzanie regularnych audytów etycznych, ocen ryzyka dla nowo wdrażanych modeli oraz monitorowanie ich działania pod kątem zgodności z przepisami i wewnętrznymi politykami (Themann, 2025).

Zarządzanie ryzykiem musi być procesem ciągłym, obejmującym cały cykl życia systemu AI – od etapu pozyskania i przygotowania danych, przez rozwój, walidację i wdrożenie modelu, aż po jego monitorowanie i ostateczne wycofanie

z użytku. Szczególną uwagę należy zwrócić na ocenę ryzyka pochodzącego z łańcucha dostaw, w tym na audyty bezpieczeństwa zewnętrznych dostawców modeli, danych i platform chmurowych (Locascio, 2023). Niezbędnym elementem jest również podnoszenie świadomości i kompetencji w zakresie bezpieczeństwa danych i zagrożeń AI wśród wszystkich pracowników. Jak pokazują statystyki, typowy cyberatak często zaczyna się od prostego błędu ludzkiego, takiego jak kliknięcie w złośliwy link. Niestety, badania wskazują na wciąż niezadowalający poziom świadomości społecznej na temat bezpiecznego korzystania z AI i ochrony danych osobowych, co czyni regularne szkolenia kluczowym elementem strategii obronnej (Khadka & Ullah, 2025).

W dobie transformacji cyfrowej napędzanej przez sztuczną inteligencję, bezpieczeństwo danych przestaje być postrzegane jako wyizolowana funkcja techniczna czy centrum kosztów. Ewoluuje ono do roli strategicznego czynnika warunkującego fundamentalną gotowość technologiczną i operacyjną całego przedsiębiorstwa. Zdolność do zapewnienia niezachwianej poufności, integralności i dostępności danych w całym, złożonym cyklu życia systemów AI jest fundamentem, bez którego niemożliwe jest zbudowanie trwałej wartości biznesowej opartej na tej technologii. Dla firm działających w Specjalnych Strefach Ekonomicznych, które z definicji aspirują do bycia krajowymi i międzynarodowymi liderami innowacji, perspektywa ta ma szczególne znaczenie. Inwestycja w solidną, wielowarstwową i proaktywną architekturę bezpieczeństwa AI nie powinna być traktowana jako obciążenie, lecz jako strategiczna lokata w kluczowe, niematerialne aktywa: zaufanie klientów i partnerów, zgodność z coraz bardziej restrykcyjnym otoczeniem prawnym oraz długoterminową odporność biznesową. W środowisku, gdzie granice między organizacjami zacierają się w ramach złożonych łańcuchów dostaw, a cyberprzestępczość staje się coraz bardziej wyrafinowana, zdolność do ochrony danych staje się synonimem zdolności do przetrwania i rozwoju. Przedsiębiorstwa, które zintegrują bezpieczeństwo z samą tkanką swojej infrastruktury AI – od potoków danych, przez procesy deweloperskie, aż po ramy zarządcze – zyskają trwałą przewagę konkurencyjną. Będą w stanie nie tylko wdrażać innowacje szybciej i efektywniej, ale także minimalizować paraliżujące ryzyka prawne, finansowe i reputacyjne. W coraz bardziej złożonym i niebezpiecznym cyfrowym świecie, solidne kompetencje w zakresie bezpieczeństwa danych stają się najważniejszym filarem, na którym można bezpiecznie budować przyszłość opartą na sztucznej inteligencji.

Strategiczne podejście organizacji do sztucznej inteligencji

Integracja AI z planowaniem biznesowym i strategicznym

Integracja sztucznej inteligencji z procesami planowania strategicznego przestała być postrzegana jako opcjonalna innowacja technologiczna. W realiach gospodarki opartej na danych stała się fundamentalnym warunkiem budowania i utrzymania trwałej przewagi konkurencyjnej. Proces ten wykracza daleko poza implementację pojedynczych narzędzi, wymuszając głęboką transformację w sposobie myślenia o strategii, podejmowaniu decyzji i tworzeniu wartości. Niniejszy podrozdział przedstawia analizę ewolucji paradygmatu strategicznego AI w świetle kluczowych teorii zarządzania, omawia ramy implementacyjne uwzględniające rolę przywództwa i kultury, identyfikuje wielowymiarowe bariery utrudniające integrację oraz dokonuje ewaluacji metod pomiaru sukcesu w dobie inteligentnej automatyzacji.

Ewolucja postrzegania sztucznej inteligencji w strategii korporacyjnej odzwierciedla rosnącą dojrzałość organizacji w rozumieniu jej transformacyjnego potencjału. W początkowej fazie adopcji dominowało podejście taktyczne, silnie zakorzenione w logice business-case. Głównymi motywatorami wdrożeń były wówczas cele operacyjne, takie jak poprawa efektywności, wskazywana przez 68% firm oraz bezpośrednia redukcja kosztów, stanowiąca priorytet dla 60% organizacji. W tym ujęciu AI była traktowana jako zaawansowane narzędzie służące do automatyzacji i optymalizacji istniejących, dobrze zdefiniowanych procesów. Obecnie obserwuje się fundamentalną zmianę tego paradygmatu. Badania wskazują, że aż 87% globalnych organizacji uznaje technologie AI za kluczowy czynnik umożliwiający osiągnięcie przewagi nad konkurencją. Sztuczna inteligencja jest

coraz częściej postrzegana jako strategiczna „dźwignia jakości, innowacji i przewagi rynkowej”, a nie wyłącznie instrument do cięć budżetowych. Ta reorientacja jest widoczna w strategicznych celach wyznaczanych przez przedsiębiorstwa. Poprawa innowacyjności i zdolności do wprowadzania nowych produktów (38%), profesjonalizacja oferty dla klientów (31%) oraz ogólne zwiększenie konkurencyjności (25%) stają się równie istotne, co tradycyjne cele operacyjne. AI staje się fundamentem transformacji w kierunku Marketingu 5.0, gdzie analiza danych w czasie rzeczywistym, predykcja zachowań i hiperpersonalizacja komunikacji stanowią nowy standard rynkowy (Krawiec, 2025).

Ta zmiana perspektywy z efektywności procesowej na efektywność strategiczną jest kluczowa i znajduje swoje uzasadnienie w teorii zasobowej, która zakłada, że trwała przewaga konkurencyjna wynika z posiadania unikalnych, trudnych do imitacji zasobów i zdolności (Utama & Sari, 2025). W tym ujęciu, zaawansowane systemy AI, zintegrowane z unikalnymi zbiorami danych firmy i specyficznymi procesami, przestają być jedynie narzędziem, a stają się strategicznym, niematerialnym aktywem (Nzama et al., 2024). Prawdziwa wartość strategiczna AI nie leży w optymalizacji status quo, ale w jego redefinicji. Firmy, które ograniczają zastosowanie AI do roli narzędzia oszczędnościowego, ryzykują strategiczną krótkowzroczność i utratę dystansu do liderów rynku. Dojrzałe organizacje rozumieją, że AI jest platformą do eksploracji nowych modeli biznesowych, identyfikacji nisz rynkowych i tworzenia unikalnej wartości dla klienta. Integracja AI ze strategią pozwala na fundamentalną zmianę w sposobie podejmowania decyzji – od reaktywnego analizowania danych historycznych do proaktywnego, predykcyjnego i preskryptywnego zarządzania. Co więcej, sztuczna inteligencja demokratyzuje dostęp do zaawansowanych narzędzi analitycznych, które niegdyś były domeną wyłącznie największych korporacji, otwierając nowe możliwości dla małych i średnich przedsiębiorstw (Costa et al., 2024).

Skuteczna integracja sztucznej inteligencji z procesami biznesowymi wymaga czegoś więcej niż tylko inwestycji w technologię; wymaga przede wszystkim posiadania spójnej i dobrze zakomunikowanej strategii AI. Jej brak jest trzecią co do wielkości barierą w implementacji, wskazywaną przez 26% firm. Problem ten jest szczególnie dotkliwy w dużych organizacjach, gdzie odsetek ten wzrasta do 24% w porównaniu z 18% w średnich przedsiębiorstwach. Efektywna strategia musi definiować konkretne, mierzalne cele, które są ściśle powiązane z nadrzędnymi celami biznesowymi organizacji. Wybór odpowiedniej strategii jest kluczowy dla skutecznego zarządzania zasobami ludzkimi i finansowymi w kontekście adopcji AI, a jej ramy powinny uwzględniać zarówno aspekty formalne, jak

i kulturę organizacyjną, która sprzyja innowacjom (Krawiec, 2025; Myszak et al., 2025).

Proces integracji strategicznej można podzielić na kilka kluczowych etapów:

- Audyt danych i zasobów - przed rozpoczęciem jakichkolwiek wdrożeń technicznych kluczowe jest przeprowadzenie solidnego audytu danych w celu oceny ich jakości, dostępności i bezpieczeństwa. Problemy z jakością lub dostępnością danych stanowią istotną barierę dla 23% firm (SAP, 2024; Krawiec, 2025);
- Faza pilotażowa - inicjowanie projektów pilotażowych pozwala na osiągnięcie szybkich zwycięstw, przetestowanie technologii w kontrolowanym środowisku i zbudowanie wewnętrznego poparcia dla dalszych inwestycji. Jednakże wiele firm napotyka trudności w efektywnym skalowaniu obiecujących wyników z fazy testowej (Agbaakin, 2025);
- Skalowanie i integracja z modelem operacyjnym - pełna integracja AI wymaga fundamentalnych zmian w modelu operacyjnym, strukturach organizacyjnych i procesach. W sektorze produkcyjnym już 48% przedsiębiorstw dokonało takich modyfikacji, co świadczy o zaawansowaniu transformacji w tej branży (Mikalef & Gupta, 2021);
- Monitoring i pomiar efektywności - niezbędne jest zdefiniowanie formalnych kluczowych wskaźników efektywności (KPI) i prowadzenie regularnego monitoringu. Obecnie stanowi to wyzwanie, gdyż zaledwie 16% firm posiada sformalizowane systemy pomiaru, a większość opiera się na analizach ad-hoc (Licea, 2025).

Sukces wdrożenia zależy również od zdolności do wykorzystania ekosystemu zewnętrznego. Krajowe i międzynarodowe strategie rozwoju AI kładą silny nacisk na ścisłą integrację świata nauki z biznesem oraz przyspieszoną komercjalizację wyników badań. Rządy i instytucje publiczne odgrywają kluczową rolę jako zlecniodawcy projektów i promotorzy gospodarki opartej na danych, tworząc ramy finansowe i prawne sprzyjające innowacjom. Sojusze strategiczne mogą wspierać akceptację i chęć wdrożenia nowych technologii, chociaż obecnie niewiele firm decyduje się na tę formę współpracy (Fenwick et al., 2024).

W kontekście dynamicznych zmian rynkowych, samo posiadanie zasobów (jak technologia AI) jest niewystarczające. Kluczowa staje się teoria zdolności dynamicznych, która podkreśla umiejętność organizacji do „ciągłego budowania, integrowania i rekonfigurowania swoich umiejętności w celu adaptacji do otoczenia i utrzymania przewagi konkurencyjnej”. Sztuczna inteligencja staje się

potężnym narzędziem wzmacniającym te zdolności. Umożliwia ona wyczuwanie nowych trendów rynkowych poprzez zaawansowaną analitykę danych, chwytanie okazji poprzez wspieranie szybkiego i trafnego podejmowania decyzji, oraz rekonfigurację zasobów i procesów poprzez inteligentną automatyzację. Sukces transformacji cyfrowej zależy od gotowości organizacji do proaktywnego przekształcania i zarządzania swoimi zasobami (Reda & Jamal, 2025). W tym procesie nieocenioną rolę odgrywa przywództwo. Badania wskazują, że przywództwo transformacyjne, które buduje zaufanie, wspólną wizję i satysfakcję z pracy, jest silnie powiązane z niższymi wskaźnikami odejść pracowników w okresach zmian, co ma kluczowe znaczenie dla utrzymania talentów w dobie transformacji cyfrowej. Liderzy muszą wykazywać się elastycznością i zdolnością do innowacji menedżerskich, ponieważ brak tych cech często prowadzi do niedopasowania między wdrażanymi narzędziami technologicznymi a celami organizacji. To właśnie kadra zarządzająca jest odpowiedzialna za tworzenie kultury organizacyjnej napędzanej innowacjami, która motywuje pracowników i redukuje opór przed zmianą. Organizacje, które łączą wdrażanie nowych narzędzi cyfrowych z aktualizacją swoich modeli zarządczych, np. poprzez wprowadzanie zwinnych zespołów czy partycypacyjnego stylu przywództwa, osiągają znacznie lepsze wyniki (Reda & Jamal, 2025). Poniższa tabela syntetycznie przedstawia fundamentalną różnicę między tradycyjnym a zintegrowanym z AI podejściem do planowania strategicznego (tab. 4).

Tabela 4. Porównanie tradycyjnego i wspomaganego przez AI planowania strategicznego

Kryterium	Tradycyjne planowanie strategiczne	Planowanie strategiczne zintegrowane z AI
Analiza danych	Oparta na danych historycznych, często fragmentaryczna i opóźniona.	Analiza predykcyjna i preskryptywna w czasie rzeczywistym, przetwarzanie ogromnych zbiorów danych.
Szybkość decyzji	Cykliczna (np. roczne przeglądy strategii), reaktywna na zmiany rynkowe.	Ciągła, proaktywna, ze wsparciem lub pełną automatyzacją decyzji operacyjnych.
Prognozowanie	Ekstrapolacja trendów historycznych, modele o ograniczonej złożoności.	Zaawansowane symulacje scenariuszowe, modelowanie złożonych, nieliniowych systemów.
Personalizacja	Szeroka segmentacja rynku oparta na danych demograficznych.	Hiperpersonalizacja oferty, komunikacji i doświadczenia klienta (CX) w skali 1:1.

Kryterium	Tradycyjne planowanie strategiczne	Planowanie strategiczne zintegrowane z AI
Innowacyjność	Zależna od ludzkiej kreatywności i analizy rynku; proces często powolny.	Wspierana przez generowanie nowych insightów, identyfikację niszy rynkowych i automatyzację prac B+R.

Źródło: Opracowanie własne na podstawie: Jamarani, E., Haddadi, M., Sarvizadeh, A., Kashani, M. H., Akbari, R., & Moradi, H. (2024). Big data and predictive analytics: A systematic review of applications. *Artificial Intelligence Review*, 57, 176; Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83; Krawiec, P. (2025). Korzyści z wdrożenia AI w biznesie: efektywność, decyzje i nowe źródła przychodów. *Mobile Trends*. <https://mobiletrends.pl/korzyści-z-wdrożenia-ai-w-biznesie-efektywnosc-decyzje-i-nowe-zrodla-przychodow/>; Xie, Y., Chen, Y., Wei, Q., & Yin, H. (2024). A hybrid deep learning approach to improve real-time effluent quality prediction in wastewater treatment plant. *Water Research*, 250, 121092; Vogel, A. T., Vogel, J., Watchravesringkan, K., Cook, S. C., Beasley, J., Croom, R., Peterson, D., & Finkelstein, J. (2023). Design piracy: An interdisciplinary investigation into competitive industrial behavior. *Journal of Business Research*, 164, 113946; Huang, H., Wang, F., Song, M., Baležentis, T., & Streimikiene, D. (2021). Green innovations for sustainable development of China: Analysis based on the nested spatial panel models. *Technology in Society*, 65, 101593.

Pomimo rosnącej świadomości potencjału AI, przedsiębiorstwa napotykają na szereg złożonych barier, które hamują proces jej strategicznej integracji. Bariery te można podzielić na dwie główne kategorie: „twarde”, związane z zasobami i technologią, oraz „miękkie”, dotyczące strategii, kultury i regulacji. Wśród barier „twardych” na pierwszym miejscu znajdują się kwestie finansowe – wysokie koszty wdrożenia i utrzymania systemów AI są główną przeszkodą dla 39% firm. Drugą kluczową barierą jest brak specjalistycznych kompetencji (29%) oraz dedykowanego zespołu AI (14%). Problem ten jest pogłębiany przez tzw. „lukę kompetencyjną AI”, czyli rosnącą rozbieżność między umiejętnościami dostępnymi na rynku pracy a zapotrzebowaniem dynamicznie rozwijającej się technologii. Trzecim elementem są dane – problemy z ich jakością, dostępnością lub integracją stanowią wyzwanie dla 23% organizacji (Krawiec, 2025).

Bariery „miękkie” mają charakter strategiczny i kulturowy. Jak wspomniano, brak jasnej strategii AI (26%) jest jednym z głównych hamulców. Równie istotny jest opór kulturowy wewnątrz organizacji. Wynika on często z obaw pracowników o utratę pracy (19% badanych) oraz sceptycyzmu co do realnych korzyści (21% badanych). Kluczową barierą staje się zaufanie, które jest kształtowane przez czynniki techniczne (dokładność, transparentność, wyjaśnialność), organizacyjne (wsparcie liderów, ład korporacyjny, zarządzanie ryzykiem) oraz związane z interakcją człowiek-maszyna (kontrola użytkownika, postrzegana uczciwość).

Rosnące znaczenie mają również ograniczenia regulacyjne i prawne (19%) oraz kwestie etyki algorytmicznej, w tym ryzyko stronnictwa i dyskryminacji w algorytmach (Krawiec, 2025).

Istotne jest zrozumienie, że bariery te nie są odizolowanymi problemami, lecz tworzą system naczyń połączonych, swoistą „spiralę bezwładności”. Brak spójnej strategii (bariera miękka) bezpośrednio utrudnia uzasadnienie wysokich kosztów inwestycji (bariera twarda) oraz zdefiniowanie potrzebnych kompetencji. Z kolei brak komunikacji strategicznej ze strony zarządu wzmacnia opór kulturowy i obawy pracowników (bariery miękkie). Próba przełamania tylko jednej z tych barier, np. poprzez zatrudnienie zespołu data scientist, bez jednoczesnego adresowania kwestii strategii, kultury i zarządzania zmianą, jest skazana na niepowodzenie. Problem integracji AI jest w swej istocie problemem przywództwa strategicznego, a nie wyłącznie wyzwaniem technologicznym (Katragadda, 2025).

W sektorze produkcyjnym pojawiają się dodatkowe, specyficzne wyzwania. Kluczowe staje się zapewnienie bezpieczeństwa nie tylko systemów IT, ale również technologii operacyjnych, które bezpośrednio sterują procesami produkcyjnymi. Naruszenie bezpieczeństwa systemu AI zarządzającego pracą robotów przemysłowych może prowadzić do natychmiastowego wstrzymania całej produkcji. Mimo że 94% firm deklaruje weryfikację kwestii bezpieczeństwa przy wdrażaniu AI, realne inwestycje w zabezpieczenie tych systemów podejmuje zaledwie 34%. Ten wyraźny rozdźwięk między deklaracjami a działaniami stanowi krytyczne ryzyko strategiczne dla całego sektora (Olak, 2025).

Jedną z istotnych barier w adopcji sztucznej inteligencji są trudności w precyzyjnej ocenie zwrotu z inwestycji (ROI), na co wskazuje 13% firm. Problem ten wynika z faktu, że wiele korzyści płynących z integracji AI ma charakter jakościowy i długoterminowy – np. poprawa jakości podejmowanych decyzji, wzrost zdolności innowacyjnych czy wzmocnienie satysfakcji klienta. Te niematerialne aktywa trudno jest skwantyfikować za pomocą tradycyjnych modeli finansowych. W rezultacie, zaledwie 16% organizacji posiada formalnie zdefiniowane wskaźniki KPI i prowadzi regularny monitoring efektywności AI, podczas gdy większość opiera się na pomiarach prowadzonych ad-hoc lub dopiero planuje wdrożenie odpowiednich metryk (Krawiec, 2025a; Krawiec, 2025b).

Obserwuje się jednak wyraźną ewolucję w podejściu do pomiaru sukcesu. Firmy odchodzą od wskaźników skoncentrowanych wyłącznie na kosztach i efektywności na rzecz bardziej zniuansowanych metryk, które uwzględniają jakość i doświadczenie kluczowych interesariuszy. Chociaż poprawa efektywności procesów (46%) i redukcja kosztów (44%) pozostają ważnymi miernikami, coraz

większe znaczenie zyskują wskaźniki takie jak poprawa doświadczenia pracowników (EX), wskazywana przez 44% firm, wzrost satysfakcji klientów (CX) (40%) oraz szybkość podejmowania decyzji (22%). Ta zmiana jest szczególnie widoczna w miarę dojrzewania wdrożeń AI w organizacji: na etapie planowania dominują cele kosztowe, ale w fazie regularnego użycia kluczowa staje się szybkość decyzji, doceniana przez 83% dojrzałych użytkowników AI (Krawiec, 2025b).

Ta zmiana w systemie pomiaru jest bezpośrednim odzwierciedleniem dojrzałości strategicznej organizacji w zakresie AI i wpisuje się w logikę nowoczesnych ram pomiaru efektywności, takich jak Zrównoważona Karta Wyników. Model ten, odchodząc od czysto finansowej oceny, integruje perspektywę klienta, procesów wewnętrznych oraz uczenia się i rozwoju. W kontekście AI, perspektywa klienta może być mierzona wskaźnikami CX, perspektywa procesów – metrykami efektywności i automatyzacji, a perspektywa uczenia się i rozwoju – zdolnością do innowacji i adaptacji nowych kompetencji przez pracowników. Takie holistyczne podejście pozwala uchwycić pełne spektrum wartości tworzonej przez AI, daleko wykraczające poza proste oszczędności kosztowe (Marques & Oliveira, 2024). Rodzaj stosowanych wskaźników KPI może służyć jako narzędzie diagnostyczne do oceny, czy sztuczna inteligencja jest traktowana jako narzędzie taktyczne, czy jako strategiczny filar działalności. Przedsiębiorstwa na niskim poziomie dojrzałości koncentrują się na KPI operacyjnych, takich jak redukcja kosztów obsługi klienta dzięki chatbotom. Z kolei organizacje dojrzałe strategicznie wdrażają holistyczne wskaźniki, które mierzą wpływ AI na kluczowe filary przewagi konkurencyjnej, np. wskaźnik retencji klientów, czas wprowadzania nowego produktu na rynek czy poziom zaangażowania pracowników. W tym ujęciu system pomiaru efektywności przestaje być jedynie narzędziem kontrolnym, a staje się lustrem, w którym odbija się głębokość i jakość strategicznej integracji AI z modelem biznesowym firmy (Damnjanović et al., 2025).

Integracja sztucznej inteligencji z planowaniem biznesowym i strategicznym jest złożonym procesem transformacji, który wykracza daleko poza implementację technologii. Obejmuje on fundamentalne zmiany w strategii, kulturze organizacyjnej, wymaganych kompetencjach oraz modelach przywództwa. Jak pokazują teorie zasobowa i zdolności dynamicznych, AI może stać się unikalnym, trudnym do skopiowania zasobem, który pozwala firmom nie tylko optymalizować bieżącą działalność, ale przede wszystkim dynamicznie adaptować się do zmian w otoczeniu. Sukces w tej dziedzinie nie zależy od samej technologii, ale od zdolności organizacji do holistycznego zarządzania zmianą, wspieranego przez transformacyjne przywództwo i kulturę opartą na danych. Przedsiębiorstwa,

które potrafią przejść od myślenia o AI w kategoriach optymalizacji kosztów do myślenia w kategoriach tworzenia nowej wartości, zyskują trwałą przewagę konkurencyjną. W przyszłości integracja ta będzie się pogłębiać, zwłaszcza w obszarze zaawansowanych analiz predykcyjnych i tworzenia autonomicznych systemów decyzyjnych, co postawi przed menedżerami i strategami kolejne, jeszcze bardziej złożone wyzwania, szczególnie w kontekście etyki, transparentności i zarządzania ryzykiem (Robertson et al., 2025).

Rola kierownictwa i zarządów w inwestycjach w AI

Adopcja sztucznej inteligencji w przedsiębiorstwie stanowi fundamentalną transformację strategiczną, a nie jedynie technologiczną modernizację. Wymaga to bezpośredniego, odgórnego zaangażowania liderów, którzy muszą pokierować organizacją w przejściu od marginalnych, izolowanych eksperymentów do głębokiej i strategicznej integracji AI z kluczowymi procesami tworzącymi wartość. Proces ten nie jest prostym ćwiczeniem inżynierskim, lecz złożonym wyzwaniem behawioralnym, zakorzenionym w zasadach zarządzania zmianą, które mogą być skutecznie nawigowane wyłącznie przez kadrę zarządzającą. Transformacja technologiczna, z definicji, musi rozpoczynać się na najwyższych szczeblach organizacji, gdzie liderzy aktywnie angażują się w rozwój technologiczny przedsiębiorstwa i wyznaczają strategiczne kierunki (Tursunbayeva & Gal, 2024).

Specjalne Strefy Ekonomiczne stanowią unikalny ekosystem o wysokim potencjale innowacyjnym. Funkcjonują jako „silniki wzrostu gospodarczego” i „magnesy dla firm technologicznych”, oferując uproszczone regulacje, zaawansowaną infrastrukturę cyfrową oraz ułatwiony dostęp do wykwalifikowanych kadr i kapitału. Strefy te, poprzez subsydia i efekt aglomeracji, tworzą środowisko sprzyjające transferowi wiedzy i technologii, co obniża bariery wejścia dla zaawansowanych rozwiązań, takich jak AI. Jednakże same te sprzyjające warunki nie gwarantują sukcesu. Stanowią one jedynie uśpiony potencjał, który wymaga aktywacji poprzez świadome i zdecydowane działania przywódcze (Zeng, 2021). W tym kontekście wyłania się koncepcja symbiozy między przywództwem a ekosystemem innowacji. Wrodzone zalety SSE – takie jak zachęty podatkowe, dostęp do infrastruktury badawczo-rozwojowej czy klastry technologiczne – tworzą żyzny grunt dla innowacji, ale bez proaktywnego kierownictwa pozostają niewykorzystane. To zarząd i rada nadzorcza podejmują kluczowe decyzje o alokacji kapitału, restrukturyzacji procesów i zarządzaniu ludzkim wymiarem transformacji. Zatem firma działająca w SSE, lecz pozbawiona wizjonerskiego

przywództwa, nie zdoła w pełni skapitalizować atutów strefy. Z kolei organizacja z odważnym i kompetentnym kierownictwem może wykorzystać te same korzyści do osiągnięcia wykładniczego wzrostu. SSE dostarcza zasobów i sprzyjającego otoczenia, ale to przywództwo definiuje strategię i zapewnia jej skuteczną egzekucję. W ten sposób zwrot z jakości przywództwa jest znacznie wyższy w ramach SSE niż poza nią, co czyni rolę zarządów kluczowym czynnikiem determinującym sukces w erze AI. Wprowadzenie sztucznej inteligencji do organizacji wymaga ewolucji tradycyjnych modeli zarządzania i narodzin nowego paradygmatu, określanego w literaturze jako „przywództwo AI” (AI Leadership). Nie jest to jedynie rozszerzenie kompetencji cyfrowych, ale fundamentalna zmiana polegająca na stworzeniu symbiotycznej relacji między ludzką intuicją a inteligencją maszynową. Przywództwo AI wymaga od liderów rozwinięcia zupełnie nowych, hybrydowych kompetencji, obejmujących zarówno zdolności techniczne i adaptacyjne, jak i transformacyjne. Do kluczowych umiejętności należą: świadomość technologiczna i algorytmiczna, nadzór etyczny oraz zdolność do efektywnego kierowania zespołami hybrydowymi, złożonymi z ludzi i systemów AI. Liderzy nie muszą być ekspertami w programowaniu modeli uczenia maszynowego, ale muszą rozumieć zasady ich działania, potencjalne możliwości i, co równie ważne, ograniczenia (Mahu et al., 2025).

Niezbędny zestaw kompetencji dla lidera ery AI wykracza daleko poza techniczną biegłość. Obejmuje on przede wszystkim myślenie krytyczne i odporność poznawczą, czyli zdolność do kwestionowania wyników generowanych przez AI, weryfikacji źródeł i identyfikacji ukrytych uprzedzeń w danych. W środowisku, gdzie technologia automatyzuje coraz więcej zadań analitycznych, rośnie znaczenie inteligencji emocjonalnej i empatii. Autentyczna obecność lidera, jego zdolność do aktywnego słuchania, reagowania na obawy zespołu i budowania zaufania stają się kluczowe w obszarach, w których człowiek pozostaje niezastąpiony. Ponadto, dynamiczne tempo zmian technologicznych wymusza postawę uczenia się przez całe życie. Liderzy, którzy sami inwestują w rozwój kompetencji, dają przykład i inspirują całą organizację do budowania kultury adaptacyjnej (Tarnas, 2025).

W procesie transformacji kluczową rolę odgrywa zjawisko „symbolizacji AI przez lidera”, które polega na widocznym demonstrowaniu przez kierownictwo akceptacji i wsparcia dla sztucznej inteligencji poprzez swoje zachowania i decyzje. Takie działania, jak wdrażanie projektów opartych na AI czy publiczne korzystanie z narzędzi analitycznych, działają jak potężny sygnał dla całej organizacji, komunikując strategiczne znaczenie nowej technologii. Zgodnie z teorią sygnału,

takie zachowania liderów redukują asymetrię informacyjną i łagodzą obawy pracowników, wpływając pozytywnie na ich gotowość do zmiany (Hu et al., 2025). Należy jednak podkreślić, że symboliczne działania liderów mają dwoistą naturę. Z jednej strony, wysyłając jasny sygnał o strategicznym priorytecie, mogą motywować pracowników i zachęcać do proaktywnego kształtowania swoich ról w nowej rzeczywistości. Z drugiej strony, te same działania, jeśli nie są poparte odpowiednimi zasobami i komunikacją, mogą nasilać postrzegane zagrożenia, prowadzić do oporu i lęku przed utratą pracy. Kluczowym czynnikiem modyfikującym ten efekt jest postrzegane wsparcie organizacyjne. Kiedy symbolicznym gestom lidera towarzyszą konkretne inwestycje w rozwój pracowników – szkolenia, programy reskillingu, transparentna komunikacja na temat przyszłości ról zawodowych – pracownicy interpretują zmianę jako wyzwanie, któremu mogą sprostać. W próżni informacyjnej i bez realnego wsparcia, ten sam sygnał jest odbierany jako bezpośrednie zagrożenie dla ich stabilności zawodowej. Oznacza to, że obowiązkiem zarządu jest nie tylko promowanie AI, ale przede wszystkim zapewnienie, że budżety transformacyjne holistycznie łączą koszty wdrożenia technologii z kosztami rozwoju kapitału ludzkiego. Symbol bez wsparcia staje się bowiem źródłem niepokoju i może przynieść efekt przeciwny do zamierzonego (Hu et al., 2025).

Rada nadzorcza i zarząd pełnią rolę głównych architektów i strażników ładu korporacyjnego dla sztucznej inteligencji (AI Governance). Nie jest to zadanie, które można delegować wyłącznie do działu IT; stanowi ono fundamentalny element obowiązku powierniczego najwyższego kierownictwa. AI Governance definiuje się jako zbiór procesów, struktur i standardów, które przekładają abstrakcyjne zasady etyczne na konkretne, praktyczne działania na poziomie organizacji, zapewniając, że systemy AI działają w sposób odpowiedzialny, transparentny i zgodny z prawem oraz wartościami firmy. Ustanowienie solidnych ram zarządzania AI stało się strategiczną koniecznością dla nowoczesnych przedsiębiorstw, a odpowiedzialność za ich wdrożenie spoczywa bezpośrednio na zarządzie. Do kluczowych zadań kierownictwa należy m.in. powoływanie komitetów ds. etyki AI, przyjmowanie wewnętrznych wytycznych etycznych oraz zapewnienie międzyfunkcjonalnej współpracy w celu wypracowania spójnego podejścia do wdrażania AI (Sharma et al., 2025).

Efektywny nadzór na poziomie zarządu można ustrukturyzować wokół czterech kluczowych wymiarów: strategii i konkurencyjności, alokacji kapitału, zarządzania ryzykiem oraz kompetencji technologicznych. Każdy z tych obszarów wymaga od rady nadzorczej i zarządu podjęcia konkretnych działań, aby zapewnić,

że inwestycje w AI przynoszą wartość, jednocześnie minimalizując związane z nimi zagrożenia. Operacjonalizacja ładu korporacyjnego wymaga praktycznych kroków, takich jak precyzyjne definiowanie przypadków użycia dla każdego systemu AI, przypisanie jednoznacznej ludzkiej odpowiedzialności za jego działanie i wyniki, zapewnienie wyjaśnialności podejmowanych przez algorytmy decyzji oraz systematyczne testowanie pod kątem uprzedzeń i potencjalnych szkód. Poniższa tabela syntetyzuje te kluczowe wymiary, oferując praktyczny przewodnik dla organów zarządczych (Sundararajan, 2025) (tab. 5).

Fundamentalne znaczenie ma zrozumienie, że ład korporacyjny dla AI nie może być statycznym dokumentem czy jednorazowo wdrożoną polityką. Musi on funkcjonować jako dynamiczny, adaptacyjny system. Tradycyjne modele zarządzania często opierają się na zgodności ze stabilnymi, dobrze zdefiniowanymi regulacjami. Systemy AI, zwłaszcza te oparte na uczeniu maszynowym, są z natury niestacyjne – ich wydajność może ulegać zmianom w miarę napływu nowych danych, co jest zjawiskiem znanym jako „dryf modelu”. Jednocześnie zewnętrzne otoczenie prawne i regulacyjne dotyczące AI jest niezwykle dynamiczne i często fragmentaryczne. W rezultacie, ramy zarządcze stworzone dzisiaj, jeśli pozostaną niezmienione, szybko staną się przestarzałe i nieadekwatne, tworząc iluzję bezpieczeństwa. Wymaga to zmiany paradygmatu w myśleniu zarządu: od postrzegania swojej roli jako „strażnika” zgodności do roli „zarządcy” adaptacyjnego systemu. System ten musi zawierać wbudowane mechanizmy ciągłego monitorowania, regularnych audytów modeli, kanałów informacji zwrotnej w czasie rzeczywistym oraz jasnych protokołów aktualizacji zasad i zabezpieczeń. Innymi słowy, sam system zarządzania musi uczyć się i ewoluować, odzwierciedlając naturę technologii, którą nadzoruje (Eisenberg et al., 2025).

Ocena inwestycji w sztuczną inteligencję stanowi jedno z największych wyzwań dla zarządów, ponieważ tradycyjne modele zwrotu z inwestycji (ROI) okazują się nieadekwatne i mogą prowadzić do błędnych decyzji strategicznych. Obliczanie ROI dla projektów AI jest złożone i nie poddaje się uniwersalnym formułom, głównie ze względu na szeroki zakres technologii i zastosowań objętych tym terminem. Badania pokazują, że firmy na wczesnych etapach adopcji AI często odnotowują niski lub zerowy zwrot z inwestycji (średnio 1.3%), a znaczące korzyści finansowe (średnio 4.3% dla liderów) materializują się dopiero po osiągnięciu odpowiedniej skali wdrożeń w całej organizacji. Co więcej, w przypadku najnowszych technologii, takich jak generatywna AI, stworzenie wiarygodnego modelu ROI jest na obecnym etapie rozwoju praktycznie niemożliwe z powodu braku danych i dojrzałości rynku (Celi & Miles, 2020).

Tabela 5. Kluczowe Wymiary Nadzoru Zarządu nad AI

Wymiar Nadzoru	Kluczowe Obowiązki i Pytania	Rekomendowane Działania
Strategia i konkurencyjność	- Czy AI jest integralną częścią strategii biznesowej? - Jak AI wpływa na nasz model biznesowy i pozycję konkurencyjną?	- Włączenie AI jako stałego punktu do corocznego spotkania strategicznego zarządu. - Analiza, jak AI może tworzyć nowe strumienie przychodów lub modele biznesowe.
Alokacja kapitału	- Jakie są kryteria zatwierdzania inwestycji w AI? - Czy budżet uwzględnia zarówno wdrożenie, jak i utrzymanie oraz rozwój kompetencji?	- Promowanie eksperymentowania i projektów pilotażowych. - Zabezpieczenie w rocznym budżecie inwestycji w platformy, narzędzia i partnerstwa zewnętrzne.
Zarządzanie ryzykiem AI	- Czy ryzyka związane z AI (etyczne, reputacyjne, prawne) są zintegrowane z ramami ERM? - Kto w organizacji jest ostatecznie odpowiedzialny za decyzje podejmowane przez AI?	- Traktowanie nadzoru nad ryzykiem AI jako kluczowego obowiązku zarządu. - Regularne przeglądy przypadków naruszeń etycznych w branży. - Wdrożenie procesu oceny wpływu AI.
Kompetencje technologiczne	- Czy zarząd i kadra kierownicza posiadają wystarczającą wiedzę (AI literacy), aby podejmować świadome decyzje? - Czy kompetencje AI są uwzględniane w planowaniu sukcesji?	- Organizowanie szkoleń i warsztatów dla członków zarządu z udziałem ekspertów. - Ustanowienie komitetu ds. technologii i innowacji. - Ocena gotowości kadry zarządzającej do realizacji agendy AI.

Źródło: Opracowanie własne na podstawie: Tjondronegoro, D. (2024). Strategic AI Governance: Insights from Leading Nations. arXiv. <https://arxiv.org/abs/2410.01819>; Okunola, O., Ahsun, M., & Blace, D. (2025). Measuring the ROI of AI-Driven Workforce Transformation Initiatives. SSRN Electronic Journal. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5225996; Pandey, A. C., Gupta, S., & Chhajer, N. (2021). ROI of AI: Effectiveness and measurement. International Journal of Advanced Research in Computer and Communication Engineering, 10(6), 749–755; Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. Nature Machine Intelligence, 1(2); Khairullah, H., Harris, M., Hadi, F., Sandhu, A., Ahmad, S., & Alshara, M. (2025). Implementing artificial intelligence in academic and administrative processes through responsible strategic leadership in the higher education institutions, 10, 1548104.

Typowe błędy popełniane przy kalkulacji ROI obejmują niedoszacowanie niepewności co do przyszłych korzyści, ignorowanie kosztów potencjalnych błędów algorytmicznych oraz nieuwzględnianie zjawiska degradacji wydajności modeli w czasie, co wymaga stałych nakładów na ich utrzymanie. W odpowiedzi na te ograniczenia, w literaturze naukowej proponuje się bardziej zaawansowane, holistyczne ramy oceny inwestycji. Jednym z takich modeli jest „holistyczny zwrot z etyki” (Holistic ROE), który rozszerza tradycyjne pojęcie ROI o dodatkowe wymiary. Model HROE identyfikuje trzy rodzaje zwrotu z inwestycji w technologii i związane z nimi ramy etyczne:

1. Bezpośrednie zwroty ekonomiczne - mieralne oszczędności kosztów, wzrost przychodów lub poprawa wydajności operacyjnej;
2. Niematerialne zwroty (reputacyjne) - wzrost zaufania klientów, wzmocnienie marki pracodawcy, poprawa relacji z regulatorami i interesariuszami, wynikające z odpowiedzialnego wdrażania AI;
3. Opcje realne (strategiczne) - nabycie nowych zdolności organizacyjnych, wiedzy i elastyczności, które otwierają drogę do przyszłych, potencjalnie bardzo zyskownych, możliwości biznesowych, niedostępnych dla konkurencji (Bevilacqua et al., 2023).

Koncepcja „opcji realnych”, zaczerpnięta z teorii finansów, jest szczególnie istotna w kontekście decyzji zarządczych dotyczących AI. Uzasadnia ona początkowe inwestycje w nowatorskie projekty o niepewnym, natychmiastowym zwrocie, postrzegając je nie jako koszt, lecz jako nabycie strategicznej opcji na przyszłość. Ten sposób myślenia pozwala rozwiązać paradoks, przed którym stają zarządy: aby osiągnąć status lidera o wysokim ROI, firma musi najpierw przejść przez etap początkujący o niskim ROI. Tradycyjna analiza finansowa, oparta na zdyskontowanych przepływach pieniężnych, prawdopodobnie odrzuciłaby taką inwestycję. Model opcji realnych zmienia perspektywę. Początkowy projekt AI – na przykład budowa infrastruktury danych czy przeszkolenie pierwszego modelu – jest analogiczny do zakupu opcji kupna. Inwestycja (koszt „premii”) jest relatywnie niewielka, a natychmiastowy zysk może być zerowy. Jednakże ta inwestycja daje firmie prawo, ale nie obowiązek, do podjęcia znacznie większych, bardziej dochodowych działań w przyszłości, gdy technologia dojrzeje lub pojawi się odpowiednia okazja rynkowa. Firma, która nie poniosła tego początkowego kosztu, nie posiada danych, umiejętności ani infrastruktury, aby szybko zareagować – jej „opcja” wygasła. Wymaga to od zarządu ewolucji ram decyzyjnych. Pytanie kluczowe zmienia się z „Jaka jest wartość bieżąca netto (NPV) tego projektu?” na

„Jaką wartość ma elastyczność i przyszłe możliwości, które zdobędziemy dzięki tej inwestycji, i czy jest ona warta poniesionych kosztów?”. To całkowicie redefiniuje dyskusję na temat inwestycji w AI na najwyższym szczeblu, przesuwając ją z płaszczyzny czysto finansowej na strategiczną.

Jednocześnie zarząd ponosi ostateczną odpowiedzialność za nadzór nad całym spektrum ryzyk związanych z AI, które mają bezpośrednie implikacje finansowe. Należą do nich ryzyka związane z prywatnością danych, cyberbezpieczeństwem, naruszeniem praw własności intelektualnej oraz zgodnością z dynamicznie zmieniającymi się regulacjami. W praktyce efektywny nadzór oznacza przyjęcie przez zarząd trytorowego podejścia:

- włączenia AI do strategii i profilu ryzyka spółki;
- oceny wpływu zastosowań AI na pracowników, klientów i innych interesariuszy;
- stałego nadzoru nad zgodnością z prawem i adekwatnością polityk, systemów informacyjnych i kontroli wewnętrznych.

Obejmuje to ustanowienie jasnych linii raportowania i metryk efektywności rozwiązań AI, regularne przeglądy jakości i legalności danych (w tym praw do danych), zarządzanie ryzykiem narzędzi stron trzecich i transakcji M&A, testy przedwdrożeniowe oraz przeglądy modeli pod kątem stronniczości, przejrzystości i wyjaśnialności. Jako ramę postępowania wskazuje się NIST AI Risk Management Framework (funkcje: govern–map–measure–manage), co implikuje przypisanie ról i odpowiedzialności oraz ścieżek eskalacji aż do poziomu zarządu. Z uwagi na dynamicznie zmieniający się krajobraz regulacyjny (m.in. podejście oparte na poziomach ryzyka w projektowanej UE AI Act), zarząd powinien zapewnić cykliczne aktualizacje na posiedzeniach oraz – gdy to istotne – rozważać stosowne ujawnienia wobec regulatorów, partnerów i akcjonariuszy” (Gregory & Austin, 2023).

W dyskursie akademickim i biznesowym wyłonił się konsensus co do kluczowych zasad etycznych, które muszą kierować całym cyklem życia systemów AI. Należą do nich:

1. Sprawiedliwość - zapewnienie, że systemy AI nie dyskryminują i nie utrwalają istniejących uprzedzeń. Wymaga to aktywnego identyfikowania i minimalizowania stronniczości w danych treningowych i algorytmach;
2. Odpowiedzialność - ustanowienie jasnych linii odpowiedzialności ludzkiej za projektowanie, wdrażanie i monitorowanie systemów AI. Każda

- decyzja podjęta przez AI musi mieć przypisaną odpowiedzialną osobę lub organizację, a cały proces musi być audytowalny;
3. Przezroczystość i wyjaśnialność - dążenie do tego, aby działanie systemów AI było zrozumiałe dla użytkowników i interesariuszy. Należy unikać tzw. „czarnych skrzynek”, których procesów decyzyjnych nie da się zinterpretować;
 4. Prywatność i bezpieczeństwo - gwarancja, że dane osobowe są przetwarzane zgodnie z prawem (np. RODO), a systemy są odporne na ataki i awarie;
 5. Niezawodność i bezpieczeństwo - zapewnienie, że systemy AI działają w sposób przewidywalny, solidny i nie stwarzają zagrożenia dla ludzi i środowiska (Sanderson et al., 2024).

Wdrażanie tych zasad w praktyce wymaga konkretnych mechanizmów zarządczych. Jednym z kluczowych podejść jest „etyka przez projektowanie”, gdzie wartości i zasady etyczne są wbudowywane w architekturę systemu AI od samego początku, a nie dodawane jako późniejsza łątka. Równie istotne jest utrzymanie znaczącego nadzoru ludzkiego nad działaniem systemów, co gwarantuje, że ostateczna odpowiedzialność za krytyczne decyzje zawsze spoczywa na człowieku (Brey & Dainow, 2024).

W kontekście zarządzania, kluczowe staje się zrozumienie zbieżności między etyką a obowiązkiem powierniczym. Tradycyjnie, obowiązki zarządu koncentrowały się na roztropności finansowej i zarządzaniu ryzykiem, podczas gdy etyka była często postrzegana jako domena odpowiedzialności społecznej biznesu. W erze AI to rozróżnienie zanika. Błąd etyczny, taki jak wdrożenie stroniczego algorytmu rekrutacyjnego (naruszenie zasady sprawiedliwości), bezpośrednio przekłada się na materialne ryzyko biznesowe – może prowadzić do procesów sądowych o dyskryminację, wysokich kar finansowych i nieodwracalnej szkody dla reputacji firmy. Podobnie, brak przejrzystości w działaniu systemu AI (naruszenie zasady transparentności) eroduje zaufanie klientów, co prowadzi do ich odejścia i spadku przychodów. Regulacje takie jak unijny Akt o sztucznej inteligencji (AI Act) kodyfikują te zasady etyczne w normy prawne, nakładając dotkliwe sankcje za ich nieprzestrzeganie. W konsekwencji, zarząd, który nie zapewnia solidnego nadzoru nad etyką AI, nie tylko zaniedbuje wartości społeczne, ale przede wszystkim nie wypełnia swojego fundamentalnego obowiązku należytej staranności i zarządzania ryzykiem korporacyjnym. Dyskusja o etyce AI na poziomie rady nadzorczej musi być zatem prowadzona nie tylko w języku wartości, ale

również w języku ryzyka, odpowiedzialności prawnej i ochrony długoterminowej wartości dla akcjonariuszy (Cowger, 2023; Themann, 2025).

Rola kierownictwa i zarządów w procesie inwestycji w sztuczną inteligencję jest wielowymiarowa i strategicznie kluczowa. Liderzy muszą pełnić funkcję nie tylko decydentów technologicznych, ale przede wszystkim strategów, architektów ładu korporacyjnego, rozważnych inwestorów i strażników etyki. Sukces wdrożenia AI zależy od zdolności kierownictwa do odejścia od postrzegania tej technologii jako narzędzia do optymalizacji kosztów i przyjęcia perspektywy fundamentalnej transformacji modelu biznesowego. Wymaga to nowego paradygmatu przywództwa, charakteryzującego się hybrydowymi kompetencjami, zdolnością do zarządzania złożonym procesem zmiany oraz świadomym kształtowaniem kultury organizacyjnej otwartej na innowacje (Bevilacqua et al., 2025). Rola najwyższego kierownictwa wobec AI ogniskuje się wokół trzech komplementarnych osi:

- Kompetencji przywódczych „AI-driven” (w tym podejmowania decyzji opartych na danych oraz zdolności prowadzenia innowacji w warunkach niepewności);
- Czynników napędzających adopcję (organizacyjnych, technologicznych i społeczno-kulturowych);
- Strategicznego wykorzystania AI jako filaru przewagi konkurencyjnej.

W tym względzie na uwagę zasługują zintegrowane ramy, w których skuteczność inicjatyw AI zależy nie tylko od cech samych liderów, lecz od ich zdolności do łączenia umiejętności przywódczych z dojrzałą infrastrukturą danych i systemami technologicznymi, kulturą organizacyjną wspierającą transformację oraz zgodnością z normami społecznymi i oczekiwaniami interesariuszy. W praktyce oznacza to m.in. ocenę gotowości kompetencyjnej zespołów, adekwatności architektury danych oraz trwałe zakorzenienie zasad etycznych i przejrzystości w projektach AI. Ramy te poszerzają klasyczną upper echelons theory przez włączenie perspektywy zasobowej (RBV) oraz instytucjonalnej (NIT), podkreślając, że decyzje zarządów kształtują nie tylko wyniki wewnętrzne, ale i legitymizację społeczną wdrożeń AI (Bevilacqua et al., 2025).

W unikalnym kontekście Specjalnych Stref Ekonomicznych rola ta nabiera szczególnego znaczenia. SSE, oferując preferencyjne warunki regulacyjne, dostęp do infrastruktury i kapitału oraz tworząc efekt aglomeracji, znacząco obniżają bariery dla adopcji zaawansowanych technologii. Tworzą one ekosystem przyjazny innowacjom, który jednak sam w sobie nie jest gwarantem sukcesu. Potencjał

strefy musi zostać aktywowany przez świadome, strategiczne i odważne decyzje zarządcze. To właśnie kierownictwo, poprzez alokację kapitału w oparciu o holistyczne modele oceny, takie jak opcje realne, oraz poprzez budowanie dynamicznych i adaptacyjnych ram ładu korporacyjnego, przekształca możliwości oferowane przez SSE w trwałą przewagę konkurencyjną (Jahid, 2025).

W tym środowisku ujawnia się zjawisko, które można określić jako „efekt wzmocnienia przywództwa”. Firma działająca poza SSE musi poświęcać znaczną część swojej energii i zasobów na pokonywanie zewnętrznych barier – regulacyjnych, infrastrukturalnych czy związanych z dostępem do talentów. Wewnątrz SSE wiele z tych barier jest zneutralizowanych lub przekształconych w czynniki sprzyjające. Oznacza to, że ta sama jednostka wysiłku zarządczego, kapitału i strategicznego skupienia przynosi w strefie znacznie większe rezultaty. Jakość przywództwa – jego wizja strategiczna, rygor w zarządzaniu i prawość etyczna – ma nieproporcjonalnie duży wpływ na końcowe wyniki w porównaniu do firm działających w mniej sprzyjającym otoczeniu. W SSE dobre przywództwo staje się doskonałe, a słabe przywództwo prowadzi do katastrofalnego zmarnowania uprzywilejowanej szansy. Dlatego zarządy firm zlokalizowanych w Specjalnych Strefach Ekonomicznych ponoszą szczególną odpowiedzialność – nie tylko za swoje organizacje, ale za efektywne wykorzystanie publicznych i prywatnych inwestycji, które tworzą ten unikalny ekosystem. Ich decyzje są ostatecznym czynnikiem determinującym, czy potencjał innowacyjny strefy zostanie przekuty w rynkowe przywództwo i zrównoważony sukces gospodarczy w erze sztucznej inteligencji.

Proaktywne podejście: eksperymentowanie, kultura innowacji i poszukiwanie nowych zastosowań AI

W dobie cyfrowej transformacji, napędzanej przez wykładniczy rozwój technologii takich jak sztuczna inteligencja, proaktywne podejście do innowacji przestaje być strategicznym wyborem, a staje się fundamentalnym warunkiem przetrwania i osiągnięcia trwałej przewagi konkurencyjnej. W dynamicznych środowiskach gospodarczych, których modelowym przykładem są Specjalne Strefy Ekonomiczne, postawa reaktywna, polegająca na adaptacji do zmian post-factum, nieuchronnie prowadzi do technologicznego zacofania i rynkowej marginalizacji. Współczesne przedsiębiorstwa, aby skutecznie konkurować, muszą fundamentalnie przekształcać swoje modele biznesowe, operacje i sposoby tworzenia wartości, w czym AI odgrywa rolę kluczową. Sztuczna inteligencja ewoluuje bowiem z roli

narzędzia pomocniczego do pozycji strategicznego partnera, zdolnego do inicjowania przełomowych innowacji i redefiniowania całych branż.

Podstawy teoretyczne dla proaktywnego podejścia do zarządzania innowacjami technologicznymi dostarcza koncepcja zdolności dynamicznych. Zgodnie z nią, trwała przewaga konkurencyjna w zmiennym otoczeniu nie wynika z posiadania statycznych zasobów, lecz ze zdolności organizacji do ciągłego wyczuwania, chwytania i rekonfiguracji swoich zasobów i kompetencji w odpowiedzi na pojawiające się szanse i zagrożenia. Proaktywna strategia wdrożenia AI jest zatem współczesną emanacją DCV. Wymaga ona od przedsiębiorstw nie tylko pasywnej adaptacji, ale aktywnego wykorzystania AI do systematycznego skanowania otoczenia rynkowego w poszukiwaniu nowych trendów i możliwości. Przykładem może być wykorzystanie przez firmę Tesla analityki predykcyjnej, napędzanej przez AI, do prognozowania potrzeb klientów i optymalizacji innowacji produktowych, co stanowi modelowy przykład fazy „wyczuwania”. Ujęcie DCV w badaniu Liang i Yubin jest dodatkowo rozszerzone o perspektywę akceptacji technologii przez klientów (TAM), co podkreśla, że efektywne wyczuwanie musi łączyć analitykę big data/AI/IoT z projektowaniem doświadczenia użytkownika. W praktyce Tesla łączy predykcyjne analizy popytu z OTA aktualizacjami, które ciągle poprawiają funkcjonalność produktu bez ingerencji w hardware, wzmacniając postrzeganą użyteczność i łatwość korzystania. W fazie chwytania rozpoznane szanse są przekuwane w modele biznesowe: sprzedaż DTC z transparentnym cennikiem i wsparciem AI w obsłudze, a także FSD w modelu subskrypcji (SaaS), co stabilizuje przychody i pogłębia relację z klientem. Z kolei rekonfiguracja obejmuje stałe przeobrażanie architektury biznesu—od producenta aut do platformy łączącej EV, autonomię, magazynowanie energii, sieć ładowania i usługi oparte na danych; temu towarzyszą AI-owe CRM, predykcyjne utrzymanie i równoważenie eksploracji z eksploatacją.” RSIS International „Empirycznie, autorzy pokazują, że trzy wymiary transformacji—smart product innovation, smart manufacturing transformation i digital-driven customer engagement—istotnie dodatnio wpływają na wyniki firmy (finansowe i niefinansowe), co wzmacnia tezę, że proaktywna strategia AI jest współczesną emanacją DCV przekładającą się na efekty biznesowe (Liang & Yubin, 2025).

Etapy – „chwytanie” i „rekonfiguracja” – realizowane są poprzez innowacje w modelach biznesowych, które stają się głównym mechanizmem transformacji w odpowiedzi na zidentyfikowane szanse. W tym kontekście, proaktywność strategiczna fundamentalnie zmienia postrzeganie AI w organizacji. Zamiast być traktowaną jako narzędzie do redukcji kosztów poprzez automatyzację istniejących

procesów – co jest typowe dla podejścia reaktywnego – sztuczna inteligencja staje się ofensywnym motorem tworzenia nowej wartości. Umożliwia generowanie przełomowych innowacji, otwieranie nieistniejących wcześniej strumieni przychodów, jak np. subskrypcyjne modele usług dla autonomicznej jazdy oraz angażowanie się w procesy współtworzenia wartości w ramach całych ekosystemów biznesowych. To przejście odzwierciedla szerszą zmianę w postrzeganiu zasobów firmy, gdzie zasoby niematerialne, takie jak inteligencja ludzka, innowacyjność i reputacja, są identyfikowane jako najbardziej prawdopodobne źródło trwałej przewagi konkurencyjnej, w przeciwieństwie do zasobów materialnych, które mają minimalny udział w ogólnej wydajności firmy (Jorzik et al., 2024). Innowacje w modelach biznesowych napędzane przez AI (AI-driven Business Model Innovation – AID-BMI) nie ograniczają się wyłącznie do transformacji operacyjnej, lecz obejmują pełne przeprojektowanie logiki tworzenia, dostarczania i przechwytywania wartości. W ujęciu autorów, AI umożliwia firmom przesunięcie z paradygmatu efektywnościowego w kierunku paradygmatu eksploracyjnego, w którym algorytmy stają się narzędziem eksperymentowania z nowymi konfiguracjami zasobów, relacji i kanałów rynkowych. Proces ten przebiega w czterech, wzajemnie powiązanych wymiarach:

- Innowacji propozycji wartości;
- Innowacji łańcucha wartości;
- Innowacji mechanizmów przychodowych;
- Innowacji relacji z ekosystemem.

Kluczową rolę odgrywa tu zdolność organizacji do generowania „dynamicznego dopasowania” pomiędzy technologią a modelem biznesowym – firmy o wysokim poziomie dojrzałości cyfrowej są w stanie szybciej rekonfigurować swoje modele w odpowiedzi na dane z rynku i algorytmy predykcyjne. AI napędza zjawisko *continuous business model innovation* – ciągłego eksperymentowania z mikroelementami modelu biznesowego, co wymaga kultury organizacyjnej sprzyjającej uczeniu się i tolerancji dla niepewności. W tym kontekście, fazy „chwytania” i „rekonfiguracji” przestają być odrębnymi etapami, a stają się dynamicznym, iteracyjnym cyklem adaptacyjnym, w którym dane, algorytmy i kreatywność ludzi współtworzą źródła nowej wartości (Jorzik et al., 2024).

Aby skutecznie realizować strategię opartą na DCV, organizacja musi rozwinąć specyficzną kompetencję, określaną jako „zdolność AI”. Nie jest to pojęcie tożsame z posiadaniem zaawansowanej technologii. Zdolność AI to holistyczna, wielowymiarowa konstrukcja obejmująca trzy kluczowe obszary: zdolność

infrastrukturalną (technologie, dane), zdolność zarządczą (wiedza, procesy, struktury) oraz zdolność biznesową (kreatywne zastosowanie AI do generowania wartości). Takie ujęcie podkreśla, że faktyczna kompetencja w zakresie AI powstaje na styku technologii, wiedzy organizacyjnej i ram instytucjonalnych, a jej istotą jest kreatywność menedżerska w wykorzystaniu potencjału AI. Co więcej, proaktywne podejście zakłada tworzenie synergicznej, symbiotycznej relacji między inteligencją ludzką (HI) a sztuczną, gdzie AI nie zastępuje, lecz uzupełnia i wzmacnia ludzkie procesy decyzyjne. Wdrożenie AI przestaje być w tym ujęciu jednorazowym projektem informatycznym z określonym początkiem i końcem. Staje się natomiast ciągłym, cyklicznym procesem strategicznym, wymagającym stałego dostosowywania budżetów, struktur zarządczych i strategii talentowych, co odzwierciedla dynamiczny charakter samej technologii i rynku. (Mikalef & Gupta, 2021).

Wdrożenie proaktywnej strategii AI jest niemożliwe bez równoległego kształtowania odpowiedniej kultury organizacyjnej. Adopcja technologii jest bowiem w swej istocie procesem społecznym, a czynniki kulturowe często okazują się decydującą barierą lub kluczowym akceleratorem sukcesu. Badania jednoznacznie wskazują, że „gotowość organizacyjna i opór kulturowy” stanowią jedne z najpoważniejszych przeszkód w implementacji AI w procesach decyzyjnych. Opór ten manifestuje się na kilku poziomach: od obaw pracowników przed utratą pracy na rzecz automatyzacji, przez lęki i niepewność związane z nową technologią, które mogą prowadzić do unikania interakcji z systemami AI i wypalenia zawodowego, aż po fundamentalne luki w kompetencjach cyfrowych i wiedzy na temat AI, które utrudniają jej zrozumienie i wykorzystanie, szczególnie wśród personelu nietechnicznego (Tsitsoula, 2025).

Skuteczna strategia musi zatem wyprzedzać te wyzwania, angażując się w proces, który można określić mianem „prewencyjnej inżynierii kulturowej”. Zamiast reagować na problemy kulturowe, gdy te już zdążą zakłócić proces wdrożenia, organizacja proaktywna inicjuje działania budujące kulturę gotowości na AI równoległe z inwestycjami technologicznymi. Kluczowe elementy takiej kultury to przywództwo, zaufanie i edukacja. Zaangażowanie i poparcie kadry zarządzającej dla odpowiedzialnych praktyk AI jest absolutnie niezbędne, aby zaufanie mogło rozprzestrzeniać się w całej strukturze firmy. Efektywne ramy zarządcze wymagają jasnego określenia odpowiedzialności liderów, tworzenia ścieżek audytu dla systemów AI oraz dedykowanych struktur zarządzania ryzykiem. Może to obejmować tworzenie nowych ról, takich jak Chief Artificial Intelligence Officer (CAIO), w celu promowania AI i nauki o danych w korporacji. Centralnym

elementem kultury gotowej na AI jest zaufanie, które nie powinno być postrzegane jako nieuchwytny odczucie, lecz jako „zdolność organizacyjna”, którą można świadomie budować i zarządzać. Fundamentem dla zaufania są ramy etyczne, transparentne struktury zarządcze oraz inwestycje w systemy „wyjaśnialnej AI”, które pozwalają użytkownikom zrozumieć logikę działania algorytmów. Zaufanie nie jest jednak stanem binarnym ani statycznym; wymaga ono „ciągłej kalibracji”. Oznacza to, że jednorazowe szkolenie czy wdrożenie transparentnego systemu nie wystarczy. Organizacja musi stworzyć mechanizmy ciągłej interakcji i pętle informacji zwrotnej, w ramach których pracownicy mogą bezpiecznie eksperymentować z AI, uczyć się na jej błędach i dynamicznie dostosowywać poziom swojego polegania na technologii. Ostatnim filarem jest systematyczna, obejmująca całą organizację edukacja w zakresie technologii, zastosowań i implikacji etycznych AI, skierowana do wszystkich szczebli – od zarządu po pracowników operacyjnych. Takie działania nie tylko podnoszą kompetencje, ale także aktywnie demistyfikują technologię i neutralizują lęki, budując fundament pod kulturę otwartą na innowacje (Meyer & Bremer, 2025; Tsitsoula, 2025; Themann, 2025).

Proces odkrywania i wdrażania nowych zastosowań AI nie jest dziełem przypadku, lecz wynikiem zdyscyplinowanej, metodycznej strategii eksperymentowania. Analiza rynkowa pokazuje wyraźną ewolucję w podejściu firm do sztucznej inteligencji: od fazy ostrożnej eksploracji i ograniczonych wdrożeń, która dominowała jeszcze w 2024 roku, do etapu, w którym AI staje się „codzienną praktyką” i „integralną częścią wewnętrznych procesów” w roku 2025, z wskaźnikiem adopcji sięgającym 97,5%. Ten gwałtowny cykl dojrzewania jest napędzany przez systematyczne eksperymentowanie, które pozwala organizacjom de-ryzykować inwestycje, przyspieszać proces uczenia się i identyfikować najbardziej wartościowe przypadki użycia. Kluczowym narzędziem w tej fazie są projekty typu „dowód słuszności koncepcji”. Służą one do testowania wykonalności i potencjału konkretnych rozwiązań AI w kontrolowanym środowisku. Przykładem mogą być cztery prototypowe agenty GIS (Geographic Information System) – LLM-Find, LLM-Geo, LLM-Cat i GIS Co-pilot – stworzone jako PoC w celu zademonstrowania zdolności AI do autonomicznego pozyskiwania danych geoprzestrzennych, przeprowadzania analiz przestrzennych i tworzenia map. Takie ukierunkowane, celowe eksperymenty pozwalają na weryfikację hipotez technologicznych i biznesowych przy minimalnym ryzyku. Równolegle, organizacje mogą wykorzystywać ewoluujące benchmarki do obiektywnej oceny i śledzenia postępów w zdolnościach systemów AI, przechodząc od statycznych testów wiedzy

do dynamicznych ewaluacji mierzących zdolność do rozwiązywania złożonych problemów w sposób zbliżony do ludzkiego (Techreviewer, 2025; Li & Yubin, 2025).

Co istotne, proces eksperymentowania pełni nie tylko funkcję techniczną, ale również kulturową. Angażując pracowników w projekty pilotażowe i PoC, organizacja tworzy bezpieczną przestrzeń do bezpośredniej interakcji z nową technologią. To praktyczne doświadczenie jest najskuteczniejszą formą budowania kompetencji cyfrowych i wiedzy o AI, bezpośrednio adresując luki zidentyfikowane jako bariera kulturowa. Sukcesy w małej skali budują zaufanie i pewność siebie w całej organizacji, łagodząc obawy i opór przed zmianą. Z tego względu, dojrzałe, proaktywne podejście do AI charakteryzuje się zmianą myślenia o celach wczesnej fazy eksperymentalnej. Głównym celem nie jest natychmiastowy zwrot z inwestycji (ROI), lecz maksymalizacja „zwrotu z nauki”. Wartością staje się zwalidowana wiedza, redukcja niepewności oraz budowa wewnętrznych zdolności, co tworzy solidny fundament pod znacznie większe, bardziej świadome i strategicznie uzasadnione inwestycje w przyszłości. Sukces tego podejścia potwierdzają dane: 82% firm, które przeszły tę ścieżkę, raportuje wzrost produktywności o co najmniej 20%, a jedna czwarta o ponad 50%. Z fazy eksperymentalnej wyłoniły się również kluczowe, obecnie powszechne zastosowania AI, takie jak generowanie kodu i dokumentacji, zautomatyzowane testowanie, analiza wymagań czy optymalizacja interfejsu użytkownika (Techreviewer, 2025).

Kulminacją proaktywnej strategii organizacyjnej, wspieranej przez kulturę innowacji i metodyczne eksperymentowanie, jest fundamentalna transformacja na poziomie modelu biznesowego. Współczesne przedsiębiorstwa coraz częściej przechodzą od koncepcji doskonalenia istniejących procesów do modelu ciągłej rekonfiguracji i tworzenia nowych źródeł wartości, w czym sztuczna inteligencja odgrywa rolę kluczowego katalizatora. AI przestaje być narzędziem pomocniczym wspierającym automatyzację, a staje się motorem innowacji organizacyjnej, umożliwiającym przeprojektowanie całych łańcuchów wartości oraz redefinicję sposobów współpracy w ramach złożonych ekosystemów gospodarczych. Sztuczna inteligencja pozwala przedsiębiorstwom tworzyć nowe modele operacyjne oparte na danych, które nie tylko usprawniają procesy, lecz także generują całkowicie nowe formy przychodów. Zdolność algorytmów uczenia maszynowego do wykrywania wzorców i zależności w danych pozwala na bardziej precyzyjne podejmowanie decyzji, prognozowanie popytu, personalizację oferty oraz automatyzację złożonych procesów logistycznych i produkcyjnych. Przekłada się to na transformację tradycyjnych struktur biznesowych w kierunku modeli

adaptacyjnych, opartych na wiedzy i współdzielonych zasobach informacyjnych. W tym kontekście AI sprzyja powstawaniu nowych koncepcji generowania i przechwytywania wartości, w których dane stanowią kluczowy zasób strategiczny. Zaawansowane systemy analityczne umożliwiają organizacjom wdrażanie złożonych modeli subskrypcyjnych, tworzenie usług cyfrowych o wysokim stopniu personalizacji oraz włączanie klientów w proces współtworzenia wartości. Dzięki temu relacje między przedsiębiorstwami a ich odbiorcami stają się bardziej partnerskie i interaktywne, co prowadzi do rozwoju nowych form współzależności w ramach cyfrowych ekosystemów biznesowych. AI napędza również mechanizm współtworzenia wartości w ekosystemach innowacji, łącząc interesariuszy z różnych sektorów w spójne sieci wymiany wiedzy, danych i zasobów. Tego rodzaju integracja sprzyja powstawaniu przełomowych innowacji technologicznych, które przekraczają granice pojedynczych organizacji i branż. Współczesne przedsiębiorstwa, wykorzystując algorytmy uczenia maszynowego oraz technologie predykcyjne, są w stanie nie tylko reagować na zmiany otoczenia, ale także je antycypować – stając się aktywnymi współtwórcami przyszłych trendów rynkowych. W sektorze produkcyjnym przełomowym przykładem zastosowania sztucznej inteligencji jest predykcyjne utrzymanie ruchu (AI-PdM), które stanowi istotne ogniwo filozofii Lean Manufacturing. Zastosowanie algorytmów predykcyjnych umożliwia bieżące monitorowanie stanu maszyn, przewidywanie awarii oraz planowanie konserwacji w sposób minimalizujący przestoje produkcyjne. W efekcie przedsiębiorstwa osiągają wyższy poziom efektywności energetycznej, redukują straty materiałowe i wydłużają żywotność sprzętu, co bezpośrednio wspiera realizację celów zrównoważonego rozwoju oraz strategii ESG. Co istotne, implementacja AI-PdM nie ogranicza się do aspektu technologicznego, lecz wymaga integracji z kulturą organizacyjną opartą na danych, współpracy interdyscyplinarnej oraz ciągłym doskonaleniu procesów. W tym modelu organizacje przechodzą od reaktywnego zarządzania do podejścia proaktywnego, w którym decyzje operacyjne podejmowane są w oparciu o prognozy i symulacje generowane przez sztuczną inteligencję. W ten sposób AI staje się nie tylko narzędziem efektywności, lecz także podstawowym mechanizmem strategicznego uczenia się organizacji. Rozwój sztucznej inteligencji prowadzi do głębokiej transformacji modeli biznesowych, które coraz częściej opierają się na danych, współpracy i adaptacyjności. AI nie tylko usprawnia procesy, lecz także umożliwia powstawanie nowych form tworzenia wartości, przyczyniając się do zrównoważonego i innowacyjnego rozwoju przedsiębiorstw w skali globalnej (Jorzik et al., 2024; Sardjono & Pratama, 2025).

Specjalne Strefy Ekonomiczne stanowią unikalny kontekst dla tego procesu, działając jednocześnie jako potężny akcelerator i źródło specyficznych wyzwań. Z jednej strony, SSE oferują „twarde” czynniki sprzyjające, takie jak zachęty polityczne (ulgi podatkowe, subsydia, preferencyjny dostęp do kredytów), które łagodzą ograniczenia finansowe i obniżają ryzyko inwestycji w nowe technologie, w tym technologie proekologiczne. Zapewniają również dostęp do najwyższej jakości infrastruktury i ułatwiają tworzenie sieci produkcyjnych, co dzięki efektom aglomeracji i transferowi wiedzy przekłada się na wzrost produktywności i sprzedaży firm. Z drugiej strony, ten sam ekosystem generuje istotne bariery. Adopcja AI w SSE jest często hamowana przez wysokie koszty, luki w infrastrukturze cyfrowej oraz, co najważniejsze, krytyczne niedobory wykwalifikowanej kadry. Pojawia się również problem „dystansu technologicznego”: lokalne firmy mogą mieć trudności z absorpcją zaawansowanej wiedzy od globalnych liderów technologicznych, co sprawia, że transfer wiedzy jest często bardziej efektywny od partnerów z innych krajów rozwijających się. Ponadto, zmodernizowane strefy charakteryzują się zaostreną konkurencją i bardziej rygorystycznymi regulacjami, np. środowiskowymi, co może prowadzić do eliminacji firm niezdolnych do szybkiej adaptacji (Wu et al., 2021; Liu et al., 2022).

Ta dwoista natura SSE prowadzi do „paradoksu Strefy”: SSE dostarcza „hardware” dla innowacji (finansowanie, infrastrukturę, sieci kontaktów), ale często brakuje im „software'u” (kompetencji, kultury organizacyjnej, zdolności do absorpcji wiedzy). Tworzy to środowisko, w którym firmy koncentrujące się wyłącznie na zakupie technologii, bez proaktywnej strategii rozwoju kapitału ludzkiego i organizacyjnego, mogą wpaść w pułapkę stagnacji. Dlatego też, aby w pełni wykorzystać potencjał SSE, strategia AI firmy musi mieć charakter ekologiczny i być silnie zorientowana na zewnątrz. Musi wykraczać poza granice samej organizacji, aktywnie dążąc do budowania partnerstw, uczestniczenia we wspólnych projektach badawczo-rozwojowych i angażowania się w rozwój lokalnego ekosystemu talentów.

ZASTOSOWANIE SZTUCZNEJ INTELIGENCJI W MARKETINGU

Od SIM do inteligentnych ekosystemów danych AI w marketingu

Zarządzanie informacją marketingową, od dekad stanowiące fundament podejmowania strategicznych decyzji w przedsiębiorstwach, przechodzi fundamentalną transformację. Tradycyjne Systemy Informacji Marketingowej (SIM), zdefiniowane w literaturze jako ustrukturyzowane procesy gromadzenia, analizy i dystrybucji danych na potrzeby marketingowe, okazały się niewystarczające w obliczu wyzwań ery cyfrowej. Ich architektura, zorientowana głównie na przetwarzanie wewnętrznych, transakcyjnych danych i generowanie okresowych raportów, nie jest w stanie sprostać trzem kluczowym wymiarom współczesnego otoczenia informacyjnego, znanym jako Big Data: ogromnej objętości, zawrotnej prędkości oraz bezprecedensowej różnorodności danych. Eksplozja danych pochodzących z mediów społecznościowych, urządzeń Internetu Rzeczy (IoT), aplikacji mobilnych oraz niezliczonych cyfrowych punktów styku z klientem uczyniła klasyczne podejście do analityki marketingowej przestarzałym (Brewis et al., 2023).

W odpowiedzi na te wyzwania narodził się nowy paradygmat, w którym centralną rolę odgrywa sztuczna inteligencja. AI, rozumiana jako zbiór technologii zdolnych do percepcji otoczenia, rozumowania, działania i uczenia się, stała się kluczowym czynnikiem umożliwiającym nie tylko przetwarzanie ogromnych zbiorów danych, ale przede wszystkim ekstrakcję z nich wartościowej wiedzy i automatyzację złożonych procesów decyzyjnych. Ta zmiana jest czymś więcej niż tylko technologicznym ulepszeniem; stanowi ona redefinicję samej istoty systemu informacyjnego. Następuje ewolucja od statycznych, opisowych systemów

SIM do dynamicznych, predykcyjnych i autonomicznych ekosystemów danych napędzanych przez AI (Jorzik et al., 2024; Sardjono & Pratama, 2024). Ta fundamentalna zmiana przesuwą rolę informacji w organizacji. Tradycyjny SIM funkcjonował jako pasywne repozytorium – dane były wprowadzane, a następnie wykorzystywane do generowania raportów opisujących przeszłe zdarzenia. Współczesny, inteligentny system jest natomiast aktywnym silnikiem poznawczym. Nie ogranicza się on do przechowywania informacji, lecz autonomicznie generuje hipotezy, prognozuje przyszłe zachowania konsumentów, identyfikuje ukryte wzorce i rekomenduje optymalne działania marketingowe. Analityka przechodzi od pytań „co się stało?” (analityka deskryptywna) i „dlaczego się stało?” (analityka diagnostyczna), przez „co się stanie?” (analityka predykcyjna), aż do „co powinniśmy zrobić?” (analityka preskryptywna). W tym nowym modelu informacja przestaje być jedynie funkcją wsparcia, a staje się proaktywnym, strategicznym zasobem napędzającym innowacje i wzrost. Przewaga konkurencyjna nie wynika już z samego faktu posiadania danych, lecz ze stopnia zaawansowania algorytmów AI wykorzystywanych do ich interpretacji i operacjonalizacji. Dla przedsiębiorstw działających w Specjalnych Strefach Ekonomicznych, których celem jest budowanie przewagi na rynkach globalnych, oznacza to konieczność inwestowania nie tylko w infrastrukturę danych, ale przede wszystkim w zaawansowane zdolności analityczne i algorytmiczne (Lepenioti et al., 2020).

Zastosowanie sztucznej inteligencji w zarządzaniu informacją marketingową nie jest monolityczne. Jej rola i stopień zaawansowania zmieniają się na poszczególnych etapach cyklu życia informacji – od pozyskiwania surowych danych po automatyzację działań. Aby w pełni zrozumieć ten proces, można posłużyć się ramą teoretyczną zaproponowaną przez Huang i Rust (2020), która wyróżnia cztery poziomy inteligencji AI: mechaniczną, analityczną, intuicyjną i empatyczną. Ta klasyfikacja pozwala na usystematyzowanie i ocenę dojrzałości technologicznej przedsiębiorstwa w zakresie wykorzystania AI w marketingu (Huang & Rust, 2020).

Etapem pierwszym jest gromadzenie i integracja danych (inteligencja mechaniczna). Na tym fundamentalnym etapie AI pełni przede wszystkim funkcję mechaniczną, automatyzując zadania związane z gromadzeniem i integracją danych na masową skalę. Systemy AI są w stanie w czasie rzeczywistym agregować dane ustrukturyzowane, pochodzące z wewnętrznych systemów, takich jak CRM (Customer Relationship Management) czy ERP (Enterprise Resource Planning), z ogromną ilością danych nieustrukturyzowanych z zewnętrznych źródeł. Obejmuje to posty w mediach społecznościowych, recenzje produktów, transkrypcje

rozmów z call center czy artykuły prasowe. Kluczową technologią jest tu przetwarzanie języka naturalnego (NLP), które pozwala na ekstrakcję znaczenia, identyfikację podmiotów oraz analizę sentymentu z danych tekstowych, przekształcając jakościowe opinie w mierzalne wskaźniki. W ten sposób AI tworzy kompletny, 360-stopniowy obraz klienta, stanowiący podstawę dla dalszych, bardziej zaawansowanych analiz (Huang & Rust, 2020).

Etap drugi polega na przetwarzaniu u analizie danych (inteligencja analityczna i intuicyjna). Po zintegrowaniu danych AI przechodzi na wyższy poziom inteligencji – analityczny. Na tym etapie algorytmy uczenia maszynowego wykraczają poza prosty opis, koncentrując się na identyfikacji wzorców i predykcji. Zaawansowane techniki klastrowania pozwalają na tworzenie mikro-segmentów klientów w oparciu o złożone, wielowymiarowe kryteria behawioralne i psychograficzne, które byłyby niemożliwe do zidentyfikowania przez analityków ludzkich. W marketingu B2B algorytmy predykcyjnego scoringu leadów oceniają prawdopodobieństwo konwersji potencjalnego klienta, pozwalając działom sprzedaży na priorytetyzację działań. Równie istotne są modele predykcji odejścia klienta, które umożliwiają proaktywne działania retencyjne. Najświeższe badania potwierdzają, że sztuczna inteligencja wchodzi obecnie w etap zaawansowanej inteligencji analitycznej, w którym hybrydowe modele ML i DL łączą klasyczne algorytmy klasyfikacji z głębokimi sieciami neuronowymi, znacząco zwiększając trafność prognoz. Wykazali, że metody takie jak XGBoost, LightGBM czy LSTM osiągają najwyższą skuteczność w predykcji churn w różnych branżach, umożliwiając segmentację klientów w czasie rzeczywistym oraz dynamiczne dostosowywanie strategii utrzymania (Imani et al., 2025). W miarę wzrostu zaawansowania, AI osiąga poziom inteligencji intuicyjnej, zdolnej do rozumienia kontekstu i subtelnych, nieliniowych zależności. Przykładem jest zaawansowana analiza sentymentu, która nie tylko klasyfikuje opinie jako pozytywne lub negatywne, ale także rozpoznaje sarkazm, ironię i niuanse emocjonalne, dostarczając głębszego wglądu w percepcję marki. Integracja uczenia zespołowego z analizą behawioralną klientów pozwala firmom na znaczną poprawę retencji oraz precyzyjne przewidywanie wartości leadów (Boozary et al., 2025). W przypadku hybrydowego podejścia do segmentacji opartego na wieloetapowym modelu: początkowe klastrowanie danych behawioralnych stanowi punkt wyjścia dla bardziej złożonych analiz predykcyjnych, co umożliwia przejście od tradycyjnych segmentów do dynamicznych mikro-segmentów aktualizowanych w czasie rzeczywistym (Kumar, 2025). Wyniki te wskazują, że etap drugi nie jest już wyłącznie fazą analizy, lecz procesem uczenia się organizacji w oparciu o dane, w którym

sztuczna inteligencja pełni funkcję „intuicyjnego analityka” – potrafiącego rozpoznać subtelne, nieliniowe zależności oraz antycypować przyszłe zachowania klientów z dokładnością niedostępną tradycyjnym metodom statystycznym.

Etap trzeci polega na generowaniu wniosków i modelowaniu predykcyjnym (inteligencja intuicyjna i analityka preskryptywna). Na tym etapie systemy AI odpowiadają na kluczowe pytanie: „co powinniśmy zrobić?”. Analityka preskryptywna wykorzystuje modele predykcyjne do generowania konkretnych rekomendacji. Jednym z najważniejszych zastosowań jest prognozowanie wartości życiowej klienta (*Customer Lifetime Value, CLV*), które pozwala na strategiczną alokację zasobów marketingowych w kierunku najbardziej dochodowych segmentów. Silniki rekomendacyjne, stanowiące rdzeń strategii personalizacji w e-commerce i platformach streamingowych, analizują zachowania użytkownika w czasie rzeczywistym, aby zasugerować mu „następną najlepszą ofertę” lub „następną najlepszą treść”, maksymalizując w ten sposób zaangażowanie i sprzedaż. Innym przykładem są algorytmy dynamicznego ustalania cen, które w oparciu o analizę popytu, cen konkurencji, pory dnia i profilu klienta, optymalizują przychody w czasie rzeczywistym. W najnowszej literaturze badacze zwracają uwagę, że etap ten stanowi przejście od analizy retrospektywnej do analityki decyzyjnej, w której sztuczna inteligencja łączy prognozowanie z działaniem. Analityka preskryptywna oparta na uczeniu maszynowym tworzy samouczące się systemy rekomendacji, zdolne nie tylko przewidywać przyszłe potrzeby klientów, lecz także generować automatyczne strategie działań w oparciu o cele biznesowe i ograniczenia operacyjne (Jorzik et al., 2024; Sardjono & Pratama, 2024; Gröger et al., 2014; Weinzierl et al., 2020; Bergman et al., 2019; Yu et al., 2023).

Ostatni, czwarty etap (automatyzacja decyzji i działań (inteligencja empatyczna) zamyka pętlę, przechodząc od wglądu do zautomatyzowanego działania. Platformy do zakupu mediów w modelu programmatic wykorzystują AI do licytacji i zakupu przestrzeni reklamowej w milisekundach, w oparciu o predykcyjne modele dotyczące wartości danego użytkownika. Systemy personalizacji dynamicznie dostosowują treść strony internetowej, ofertę produktową i komunikację e-mailową dla każdego indywidualnego odbiorcy. Granicą rozwoju jest inteligencja empatyczna, gdzie AI jest w stanie rozpoznawać i adekwatnie reagować na stany emocjonalne klienta. Inteligentne chatboty i asystenci głosowi, analizując ton głosu czy dobór słów, mogą dostosować styl komunikacji, aby lepiej zarządzać frustracją klienta lub budować pozytywną relację. Celem jest stworzenie bardziej ludzkiej, spersonalizowanej i efektywnej interakcji, co bezpośrednio przekłada

się na poprawę doświadczenia klienta (Huang & Rust, 2021; Davenport et al., 2020; Liu-Thompkins et al., 2022).

Szczególnie rewolucyjny wpływ na tym etapie ma rozwój generatywnej sztucznej inteligencji (GenAI). Modele takie jak GPT, DALL-E czy Midjourney umożliwiają automatyzację samego procesu tworzenia treści, co stanowi kwantowy skok w personalizacji. Zamiast jedynie dynamicznie dobierać gotowe elementy, systemy GenAI potrafią w czasie rzeczywistym generować unikalne teksty reklamowe, spersonalizowane e-maile, a nawet obrazy i wideo dostosowane do indywidualnego profilu odbiorcy. Ta zdolność do hiperpersonalizacji na masową skalę pozwala na tworzenie kampanii, w których każdy komunikat jest unikalnie skrojony, co prowadzi do głębszego zaangażowania i budowania silniejszych więzi emocjonalnych z marką (Patil, 2024; Kujore et al., 2025).

Proces ten ilustruje, że adopcja sztucznej inteligencji w marketingu nie jest zjawiskiem zero-jedynkowym, lecz raczej dynamicznym i wieloetapowym procesem rozwoju organizacyjnego. Przedsiębiorstwa przechodzą przez kolejne poziomy dojrzałości technologicznej i strategicznej, które odzwierciedlają stopień integracji rozwiązań AI w ich działaniach marketingowych. Początkowo wdrażane są proste automatyzacje, takie jak automatyczne wysyłki e-maili, chatboty czy systemy rekomendacyjne, które pozwalają głównie na redukcję kosztów operacyjnych i usprawnienie powtarzalnych procesów. Na kolejnym etapie organizacje zaczynają wykorzystywać analitykę predykcyjną, umożliwiającą przewidywanie zachowań klientów, optymalizację kampanii reklamowych czy lepsze zarządzanie budżetami marketingowymi — co w efekcie prowadzi do zwiększenia efektywności i trafności podejmowanych decyzji. Najwyższy poziom dojrzałości stanowią rozwiązania oparte na zaawansowanej personalizacji oraz tzw. empatii algorytmicznej, czyli zdolności systemów do rozpoznawania i reagowania na emocje, potrzeby oraz kontekst zachowań konsumenta w czasie rzeczywistym. Na tym etapie sztuczna inteligencja staje się integralnym elementem strategii doświadczenia klienta, umożliwiając tworzenie długotrwałych i spersonalizowanych relacji z odbiorcami. Pozycja firmy na tym kontinuum dojrzałości AI jest bezpośrednim wskaźnikiem jej gotowości technologicznej, elastyczności organizacyjnej oraz zdolności do adaptacji w warunkach gospodarki opartej na danych. Przedsiębiorstwa, które potrafią skutecznie integrować sztuczną inteligencję z procesami biznesowymi, zyskują nie tylko przewagę kosztową czy operacyjną, ale przede wszystkim strategiczną — budując unikalną wartość poprzez lepsze zrozumienie i zaspokajanie potrzeb swoich klientów (tab. 6).

Tabela 6. Zastosowanie technologii AI w poszczególnych etapach zarządzania informacją marketingową

Etap w cyklu życia informacji	Poziom Inteligencji AI	Kluczowe zadania marketingowe	Przykładowe technologie i algorytmy AI	Mierzalne korzyści dla przedsiębiorstwa
Gromadzenie i integracja	Mechaniczna	Agregacja danych z systemów CRM, ERP, mediów społecznościowych, IoT. Ekstrakcja danych niestrukturyzowanych.	Web scraping, integracja przez API, Przetwarzanie języka Naturalnego (NLP) do analizy tekstu.	Redukcja kosztów pozyskania danych, 360-stopniowy widok klienta, dostęp do nowych źródeł informacji.
Analiza i segmentacja	Analityczna	Identyfikacja mikro-segmentów klientów, predykcja ryzyka odejścia, scoring leadów (B2B).	Algorytmy klastrowania (np. k-średnich), drzewa decyzyjne, regresja logistyczna, sieci neuronowe.	Zwiększona precyzja targetowania, poprawa retencji klientów, optymalizacja pracy działu sprzedaży.
Generowanie wniosków i predykcja	Intuityjna	Prognostowanie wartości życiowej klienta (CLV), modelowanie atrybucji, rekomendacje produktów/treści.	Modele predykcyjne (np. LSTM dla szeregów czasowych), systemy rekomendacji (filtrowanie kolaboratywne), analiza sentymentu.	Optymalizacja budżetu marketingowego (ROI), wzrost wartości koszyka (cross/up-selling), proaktywne zarządzanie marką.
Automatyzacja i personalizacja	Analityczna / Empatyczna	Dynamiczne ustalanie cen, hiperpersonalizacja komunikacji w czasie rzeczywistym, zautomatyzowany zakup mediów.	Uczenie ze wzmocnieniem (dla cen), Generatywna AI (dla treści), systemy do programmatic advertising, inteligentne chatboty.	Maksymalizacja przychodów, wzrost zaangażowania i lojalności klientów, poprawa efektywności kampanii.

Źródło: Opracowanie własne na podstawie: Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49, 30-50; Imani, B., Joudaki, M., Belk Mohammadi, M., & Arabnia, H. R. (2025). Artificial intelligence in the context of business process management: A systematic literature review. *AI*, 7(3), 105; Boozary, P., Sheykhan, S., Ghorban Tanhaei, H., & Magazzino, C. (2025). Enhancing customer retention with machine learning: A comparative analysis of ensemble models for accurate churn prediction. *International Journal of Information Management Data Insights*, 5(1), 100331; Kumar, N. (2025). Intelligent customer segmentation: Unveiling consumer patterns with machine learning. *Journal of Umm Al-Qura University for Engineering and Architecture*, 16, 774-783; Gröger, C., Schwarz, H., & Mitschang, B. (2014). Prescriptive analytics for recommendation-based business process optimization. In W. Abramowicz & A. Kokkinaki (Eds.), *Business Information Systems. BIS 2014. Lecture Notes in Business Information Processing*, 176, 25-37; Weinzierl, S., Dunzer, S., Zillker, S., & Matzner, M. (2020). Prescriptive business process monitoring for recommending next best actions. In E. Fahland, C. Ghidini, D. Becker, & M. Dumas (Eds.), *Business Process Management Forum (BPM 2020)* (pp. 193-209). Springer. https://doi.org/10.1007/978-3-030-58638-6_12; Bergman, D., Huang, T., Brooks, P., Lodi, A., & Raghunathan, A. U. (2019). JANOS: An integrated predictive and prescriptive modeling framework. *arXiv. <https://arxiv.org/abs/1911.09461>*; Sun, Y., Liu, C., & Gao, Y. (2023). Artificial intelligence for scientific discovery and design: Past, present, and future. *Frontiers in Pharmacology*, 14, 9958434; Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2019). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24-42; Liu-Thompkins, Y., Okazaki, S., & Li, H. (2022). Artificial empathy in marketing interactions: Bridging the human-AI gap in affective and social customer experience. *Journal of the Academy of Marketing Science*, 50, 1198-1218.

Skuteczne wdrożenie sztucznej inteligencji w zarządzaniu informacją marketingową wymaga solidnego fundamentu technologicznego. Nie jest to kwestia zakupu pojedynczego oprogramowania, lecz budowy zintegrowanego ekosystemu, w którym poszczególne komponenty współdziałają w celu przekształcenia danych w działania. Ten technologiczny stos opiera się na trzech filarach: platformach danych, mocy obliczeniowej w chmurze oraz wyspecjalizowanych modelach AI. Centralnym elementem tego ekosystemu są platformy danych klienta (CDP). Stanowią one nowoczesną ewolucję tradycyjnych baz danych marketingowych. Ich kluczową funkcją jest unifikacja danych o kliencie pochodzących ze wszystkich możliwych źródeł – transakcyjnych, behawioralnych, demograficznych, online i offline – w jeden, spójny i trwały profil. Platformy te nie tylko agregują dane, ale również je normalizują i unifikują, „zszywając” informacje z różnych punktów styku (online i offline) pod unikalnym identyfikatorem klienta, co pozwala na eliminację duplikatów i zapewnia spójność danych. Dostęp do takiego ujednoliconego repozytorium danych pierwszej strony jest kluczowy w dobie wycofywania plików cookie stron trzecich i rosnących wymagań dotyczących prywatności. CDP tworzy „pojedyncze źródło prawdy” o kliencie, które staje się warstwą fundamentalną dla wszystkich aplikacji AI, zapewniając im dostęp do czystych, zintegrowanych i aktualnych danych w czasie rzeczywistym (Parker, 2024; Sell, 2025).

Drugim filarem jest chmura obliczeniowa. Trenowanie i wdrażanie zaawansowanych modeli uczenia maszynowego, zwłaszcza tych opartych na głębokim uczeniu, wymaga ogromnej mocy obliczeniowej, która jest poza zasięgiem większości przedsiębiorstw w modelu on-premise. Platformy chmurowe, takie jak Amazon Web Services (AWS), Microsoft Azure czy Google Cloud Platform, demokratyzują dostęp do tej mocy, oferując skalowalną, elastyczną i relatywnie tanią infrastrukturę. Umożliwiają one przechowywanie petabajtów danych (w tzw. jeziorach danych, data lakes) oraz dynamiczne alokowanie zasobów obliczeniowych niezbędnych do prowadzenia złożonych analiz, co jest kluczowym czynnikiem umożliwiającym wdrożenie AI na szeroką skalę (Patil, 2024; Kujore et al., 2025).

Trzeci filar to wyspecjalizowane modele i algorytmy AI. To one stanowią „mózg” operacji, analizując dane i generując wnioski. W marketingu zastosowanie znajduje szerokie spektrum modeli, w tym:

- Rekurencyjne Sieci Neuronowe (RNN) i ich warianty, jak LSTM (Long Short-Term Memory): Idealne do analizy danych sekwencyjnych,

takich jak ścieżka klienta na stronie internetowej, historia zakupów czy trendy w mediach społecznościowych. Pozwalają na modelowanie zależności czasowych i prognozowanie przyszłych zachowań;

- Konwolucyjne Sieci Neuronowe (CNN): Pierwotnie opracowane do analizy obrazów, znajdują zastosowanie w marketingu do zadań takich jak rozpoznawanie logo marki w postach w mediach społecznościowych, analiza wizualnych treści generowanych przez użytkowników czy personalizacja reklam graficznych;
- Modele Transformer (np. BERT, GPT): Stanowią przełom w przetwarzaniu języka naturalnego. Umożliwiają tworzenie zaawansowanych chatbotów prowadzących naturalne konwersacje, automatyczne generowanie spersonalizowanych tekstów marketingowych (e-maili, opisów produktów) oraz głęboką analizę semantyczną opinii klientów.

U podstaw tych modeli językowych leży architektura Transformer, która zrewolucjonizowała sposób przetwarzania danych sekwencyjnych. W odróżnieniu od wcześniejszych architektur, takich jak RNN, które analizowały dane krok po kroku, Transformatery wykorzystują mechanizm zwany „uwagą”, pozwalający na równoległe przetwarzanie całej sekwencji oraz określanie, które elementy danych są dla modelu najbardziej istotne w danym kontekście. Dzięki temu możliwe jest uchwycenie złożonych, długodystansowych zależności pomiędzy słowami czy tokenami, co znacząco zwiększa skuteczność modeli w rozumieniu znaczenia i relacji występujących w danych. Zastosowanie architektury Transformer nie ogranicza się jednak wyłącznie do języka naturalnego. Jak wskazuje AWS, modele te mogą być wykorzystywane do analizy i generowania innych rodzajów danych sekwencyjnych, takich jak muzyka, gdzie mechanizm uwagi pozwala rozpoznawać wzorce i kontynuować melodie w sposób spójny ze stylem utworu. Przykładem jest projekt AWS DeepComposer, w którym Transformer analizuje fragment melodii i generuje jej rozwinięcie, zachowując kontekst i strukturę muzyczną. AWS podkreśla również, że dzięki usługom takim jak Amazon SageMaker możliwe jest skuteczne trenowanie, dostrajanie i wdrażanie modeli Transformer w środowisku chmurowym. Integracja z biblioteką Hugging Face ułatwia implementację tego typu modeli w praktycznych zastosowaniach, takich jak analiza sentymentu, generowanie treści czy personalizacja komunikacji. W efekcie Transformatery stały się fundamentem nowoczesnych systemów sztucznej inteligencji, umożliwiając precyzyjne rozumienie kontekstu, efektywne przetwarzanie dużych ilości danych oraz tworzenie treści, które są spójne, realistyczne i dopasowane do potrzeb użytkownika. (Aws.amazon.com, 2024).

Współdziałanie tych trzech filarów tworzy symbiotyczną relację: Big Data jest bezużyteczne bez AI, które potrafi je przeanalizować, a algorytmy AI nie mogą osiągnąć swojej pełnej skuteczności bez dostępu do ogromnych, zróżnicowanych zbiorów danych. Wdrożenie i utrzymanie tych zaawansowanych modeli w środowisku produkcyjnym wymaga jednak czegoś więcej niż tylko technologii analitycznych. W odpowiedzi na to wyzwanie narodziła się dyscyplina MLOps (Machine Learning Operations), która adaptuje zasady DevOps do cyklu życia modeli uczenia maszynowego. MLOps obejmuje praktyki i narzędzia służące do automatyzacji procesów wdrażania, monitorowania, zarządzania i ponownego trenowania modeli AI. W kontekście marketingu, MLOps zapewnia, że modele predykcyjne (np. do prognozowania churnu) czy systemy personalizacji nie stają się przestarzałe, lecz są systematycznie aktualizowane w odpowiedzi na zmieniające się dane i zachowania klientów, co gwarantuje ich ciągłą skuteczność i dokładność. To właśnie interoperacyjność tych komponentów – platformy danych, chmury i modeli AI – tworzy prawdziwą wartość. Dla firm w SSE oznacza to, że skuteczna strategia AI nie polega na zakupie gotowego „rozwiązania AI”, lecz na strategicznym budowaniu elastycznego, skalowalnego stosu technologicznego, który może ewoluować wraz z rosnącymi potrzebami biznesowymi i postępowaniem technologicznym (Pahune & Akhtar, 2025).

Pomimo ogromnego potencjału, wdrożenie sztucznej inteligencji w zarządzaniu informacją marketingową wiąże się z szeregiem poważnych wyzwań o charakterze technologicznym, organizacyjnym i, co coraz ważniejsze, etycznym. Krytyczna analiza tych barier i ryzyk jest niezbędna do odpowiedzialnego i skutecznego wykorzystania tej technologii. Wśród barier organizacyjnych i technologicznych na pierwszy plan wysuwają się silosy danych. W wielu przedsiębiorstwach dane klientów są rozproszone w dziesiątkach odizolowanych systemów (CRM, systemy transakcyjne, platformy e-mailowe), co uniemożliwia stworzenie zintegrowanego profilu klienta, będącego podstawą dla AI. Integracja tych systemów jest procesem złożonym, kosztownym i czasochłonnym (Goswami, 2021). Drugim poważnym wyzwaniem jest luka kompetencyjna. Popyt na specjalistów z dziedziny data science, inżynierii uczenia maszynowego i etyki AI znacznie przewyższa podaż. Tradycyjne zespoły marketingowe, składające się z kreatywnych strategów i menedżerów marki, muszą zostać uzupełnione o nowe role lub przejść gruntowne przeszkolenie, aby móc efektywnie współpracować z systemami AI i interpretować ich wyniki. Wymaga to fundamentalnej zmiany w strukturze i kulturze organizacyjnej działu marketingu, który musi przekształcić się w interdyscyplinarne centrum analityczne (Rikala et al., 2024).

Jednakże to aspekty etyczne i regulacyjne stanowią obecnie najpoważniejsze i najbardziej złożone wyzwanie. Apetyt AI na dane wchodzi w bezpośredni konflikt z rosnącą świadomością społeczną i regulacjami dotyczącymi ochrony prywatności, takimi jak Ogólne Rozporządzenie o Ochronie Danych (RODO/GDPR). Procesy personalizacji i profilowania muszą być prowadzone w sposób transparentny, oparty na świadomej zgodzie użytkownika, co stawia przed marketerami nowe wymogi (Radanliev, 2025).

Kwestie przejrzystości, sprawiedliwości i ochrony prywatności są ze sobą głęboko powiązane i nie da się ich traktować oderwanie. W kontekście AI marketingu, gdy firma gromadzi dane, podejmuje decyzje o profilowaniu użytkowników i wykorzystuje te profile do personalizacji komunikatów, ryzyko uprzedzeń algorytmicznych i dyskryminacji staje się realne. Należy w praktyce wdrażać mechanizmy i procedury, które pomagają zminimalizować te ryzyka. Należą do nich m.in.:

- Regularne audyty etyczne i algorytmiczne, których celem jest wykrycie i skorygowanie uprzedzeń w modelach AI;
- Wdrażanie algorytmów „świadomych sprawiedliwości”, które w procesie uczenia uwzględniają kryteria równości i niedyskryminacji;
- Zapewnienie równego dostępu i reprezentacji różnych grup społecznych w zbiorach danych i zespołach rozwijających AI, co pomaga uniknąć efektu „efektu mniejszości” w modelu;
- Ustanowienie klarownych struktur odpowiedzialności i mechanizmów rozliczalności, tak aby decyzje AI — nawet jeżeli są automatyczne — mogły być audytowalne i poddane kontroli ludzkiej.

W praktyce oznacza to, że marketerzy korzystający z AI muszą wykraczać poza samo przestrzeganie prawa — powinni projektować systemy, w których uprzejmość danych, ochrona prywatności i minimalizacja szkód stają się integralną częścią procesu projektowania. Radanliev podkreśla, że przejrzystość i wyjaśnialność są kluczowe — użytkownik i regulacje muszą mieć możliwość zrozumienia, w jaki sposób dane są wykorzystywane, jakie decyzje podejmuje system AI i jakie kryteria wpływają na profilowanie i personalizację. Co istotne, autor zauważa, że różne regiony świata różnie kładą nacisk na te trzy zasady (transparentność, sprawiedliwość, prywatność) i nie ma uniwersalnego modelu etycznego akceptowanego globalnie — co stwarza wyzwania dla przedsiębiorstw działających międzynarodowo. W efekcie podejście do marketingu napędzanego przez AI musi uwzględniać nie tylko technologiczną efektywność, ale również jasno

zdefiniowane granice etyczne, nadzór, audyt oraz transparentną komunikację z użytkownikami — tak, aby systemy personalizacyjne były nie tylko skuteczne, ale i społeczne akceptowalne (Radanliev, 2025).

Kolejnym krytycznym problemem jest stronniczość algorytmiczna. Modele AI, uczone na historycznych danych, mogą nieświadomie powielać, a nawet wzmacniać istniejące w społeczeństwie uprzedzenia. Jeśli historyczne dane pokazują, że pewne grupy demograficzne rzadziej otrzymywały kredyty lub były rzadziej targetowane w kampaniach rekrutacyjnych, algorytm może nauczyć się dyskryminować te grupy, prowadząc do wykluczenia i poważnych szkód wizerunkowych. Problem ten wynika często z faktu, że dane treningowe odzwierciedlają historyczne nierówności społeczne, a algorytm, optymalizując pod kątem predykcji, uczy się tych wzorców jako pożądaných. W reklamie online może to prowadzić do systemowej dyskryminacji, gdzie pewnym grupom społecznym selektywnie odmawia się dostępu do informacji o ofertach pracy, mieszkaniach czy usługach finansowych, co utrwała istniejące podziały. Co więcej, stronniczość może ewoluować w czasie, zjawisko znane jako „dryf sprawiedliwości” pokazuje, że model uznany za sprawiedliwy w momencie wdrożenia, może z czasem zacząć generować stronnicze wyniki w odpowiedzi na zmiany w populacji lub otoczeniu. Identyfikacja i mitygacja stronniczości w danych i modelach staje się kluczowym zadaniem etycznym i biznesowym (Jonker & Rogers, 2024).

Z tym wiąże się problem przejrzystości i wyjaśnialności. Wiele zaawansowanych modeli, zwłaszcza sieci głębokiego uczenia, działa jak „czarne skrzynki”. Oznacza to, że nawet ich twórcy nie są w stanie w pełni wyjaśnić, dlaczego model podjął konkretną decyzję (np. dlaczego odrzucił wniosek kredytowy danego klienta lub zakwalifikował go jako klienta o wysokim ryzyku odejścia). Ta nieprzejrzystość stanowi poważne wyzwanie dla odpowiedzialności, zaufania i zgodności z regulacjami (np. „prawem do wyjaśnienia” w RODO). W odpowiedzi na ten problem rozwija się dziedzina Wyjaśnialnej AI (Explainable AI, XAI), której celem jest tworzenie modeli, których proces decyzyjny jest zrozumiały dla człowieka. Narzędzia XAI mają na celu dostarczenie jasnych, czytelnych dla człowieka wyjaśnień, dlaczego model podjął daną decyzję, co jest kluczowe nie tylko dla zgodności z regulacjami, ale także dla budowania zaufania użytkowników i umożliwienia marketerom świadomego nadzoru nad systemami AI. Transparentność w działaniu AI staje się fundamentem etycznego marketingu, pozwalając konsumentom zrozumieć, w jaki sposób ich dane są wykorzystywane do personalizacji (Balasubramaniam et al., 2023; Arrieta et al., 2020).

W tym kontekście, zarządzanie etyką AI przestaje być jedynie kwestią zgodności z prawem. Staje się ono potencjalnym źródłem trwałej przewagi konkurencyjnej. Firmy, które w sposób proaktywny i transparentny podchodzą do kwestii prywatności, sprawiedliwości algorytmicznej i wyjaśnialności, mogą zbudować zaufanie klientów na poziomie, który jest trudny do skopiowania dla konkurencji. Etyka przekształca się z ograniczenia w kluczowy element wartości marki. Dla przedsiębiorstw z SSE, zwłaszcza tych działających na rynkach międzynarodowych o wysokich standardach regulacyjnych, wdrożenie solidnych ram zarządzania etyką AI jest nie tylko koniecznością, ale i strategiczną inwestycją w długoterminową reputację i lojalność klientów.

Analiza ewolucji zarządzania informacją marketingową pod wpływem sztucznej inteligencji jednoznacznie wskazuje na głęboką i nieodwracalną zmianę paradygmatu. Przejście od tradycyjnych, reaktywnych Systemów Informacji Marketingowej do proaktywnych, inteligentnych ekosystemów danych jest nie tylko technologicznym krokiem naprzód, ale fundamentalną transformacją strategiczną. Korzyści płynące z tego procesu są wielowymiarowe: od bezprecedensowej hiperpersonalizacji komunikacji, która zwiększa zaangażowanie i lojalność klientów, przez znaczącą poprawę jakości decyzji biznesowych dzięki analityce predykcyjnej i preskryptywnej, aż po wzrost efektywności operacyjnej poprzez automatyzację powtarzalnych zadań analitycznych i wykonawczych (Malthouse & Copulsky, 2022; Kumar et al., 2024).

Sztuczna inteligencja powinna być postrzegana nie jako niszowe narzędzie marketingowe, ale jako technologia ogólnego przeznaczenia, która redefiniuje całą dyscyplinę marketingu, otwierając nowe możliwości tworzenia wartości i budowania relacji z klientem. W tym kontekście, zdolność przedsiębiorstwa do skutecznego wdrożenia i wykorzystania AI w zarządzaniu informacją staje się bezpośrednim i kluczowym wskaźnikiem jego ogólnej gotowości technologicznej (Triguero et al., 2023). Dla firm operujących w ramach Specjalnych Stref Ekonomicznych, które z założenia mają być centrami innowacji, technologii i działalności eksportowej, wnioski te mają szczególne znaczenie. W globalnej gospodarce, gdzie konkurencja jest coraz bardziej zintensyfikowana i oparta na danych, mistrzostwo w inteligentnym zarządzaniu informacją przestaje być opcjonalnym dodatkiem. Staje się ono fundamentalnym warunkiem przetrwania i rozwoju. Zdolność do przekształcenia surowych danych w strategiczną wiedzę i zautomatyzowane, spersonalizowane działania jest tym, co będzie odróżniać liderów rynku od pozostałych. Globalny rynek marketingu opartego na generatywnej AI, wyceniany w 2025 roku na ponad 62 miliardy dolarów, ma według prognoz

wzrosnąć do ponad 356 miliardów dolarów do roku 2030, co podkreśla skalę i tempo nadchodzącej transformacji (Hale, 2025).

Podsumowując, podróż od tradycyjnego SIM do inteligentnego ekosystemu napędzanego przez AI jest tożsama z podróżą w kierunku stania się organizacją prawdziwie opartą na danych i gotową na wyzwania przyszłości. Dla przedsiębiorstw w SSE ta transformacja nie jest jedynie kwestią optymalizacji zwrotu z inwestycji w marketing; jest to strategiczny imperatyw, warunkujący ich długoterminową konkurencyjność i pozycję na arenie międzynarodowej w coraz bardziej inteligentnym świecie.

Planowanie marketingowe wspierane przez algorytmy predykcyjne i personalizację

Współczesne planowanie marketingowe przechodzi transformację, odchodząc od tradycyjnego, retrospektywnego modelu na rzecz paradygmatu opartego na przewidywaniu i proaktywności. Zmiana ta jest napędzana przez rosnącą dostępność dużych zbiorów danych oraz postępy w dziedzinie sztucznej inteligencji, które umożliwiają przedsiębiorstwom nie tylko analizowanie przeszłych zdarzeń, ale również antycypowanie przyszłych zachowań konsumentów, trendów rynkowych i wyników biznesowych. Analityka predykcyjna, wykorzystująca algorytmy statystyczne i techniki uczenia maszynowego, stała się technologicznym fundamentem tej ewolucji, rewolucjonizując sposób, w jaki firmy tworzą i optymalizują swoje strategie. Zamiast reagować na działania klientów, organizacje zyskują zdolność do przewidywania ich potrzeb, preferencji i wzorców zakupowych, co pozwala na tworzenie spersonalizowanych doświadczeń w czasie rzeczywistym (Kumar et al., 2024; Khamoushi, 2024).

Ta zmiana paradygmatu ma głębokie implikacje strategiczne. Tradycyjny cykl planowania marketingowego, obejmujący etapy planowania, wykonania, pomiaru i wyciągania wniosków, ulega fundamentalnej modyfikacji. Analityka predykcyjna wprowadza do tego cyklu nowy, kluczowy element: antycypację. Faza wyciągania wniosków, która historycznie opierała się na analizie danych z zakończonych kampanii, teraz zasila modele predykcyjne generujące prognozy przyszłych zachowań. W rezultacie, kolejny plan strategiczny opiera się nie na tym, co było, ale na tym, co prawdopodobnie będzie. Ten proces znacząco przyspiesza podejmowanie decyzji i sprawia, że strategie stają się bardziej adekwatne i skuteczne. W perspektywie do 2030 roku, integracja AI w marketingu przestaje być postrzegana jako aspiracja, a staje się imperatywem biznesowym,

prowadzącym do w pełni zautomatyzowanych, hiperpersonalizowanych kampanii, które wyprzedzają pragnienia konsumentów, zanim te jeszcze się zmaterializują (Hunter, 2025).

W tym kontekście, zdolność do proaktywnego i antycypacyjnego marketingu staje się źródłem nowej, trwałej przewagi konkurencyjnej. Nie jest to jedynie usprawnienie operacyjne, ale fundamentalna zmiana w sposobie konkurowania na rynku. Przedsiębiorstwo, które potrafi precyzyjnie prognozować potrzeby klientów, ryzyko ich odejścia czy wartość życiową (Customer Lifetime Value - CLV), może alokować swoje zasoby ze znacznie większą precyzją niż jego konkurenci. Prowadzi to do powstania pozytywnej pętli sprzężenia zwrotnego: lepsze prognozy skutkują lepszymi doświadczeniami klientów i wyższym zwrotem z inwestycji (ROI), co z kolei generuje więcej danych, które doskonalą modele predykcyjne. Ten mechanizm tworzy solidną „fosę” konkurencyjną, trudną do pokonania dla firm o niższym stopniu dojrzałości analitycznej, co jest szczególnie istotne w dynamicznym środowisku Specjalnych Stref Ekonomicznych (Adwani, 2025).

Analityka predykcyjna jest definiowana jako zastosowanie zaawansowanych technik analitycznych, w tym algorytmów statystycznych i uczenia maszynowego, do analizy historycznych i bieżących danych w celu prognozowania przyszłych zdarzeń, trendów i zachowań (Al Khaldy et al., 2023). W marketingu jej celem jest przekształcenie surowych danych w praktyczne, przyszłościowe informacje, które umożliwiają podejmowanie bardziej świadomych decyzji strategicznych. Rdzeniem analityki predykcyjnej jest uczenie maszynowe, które pozwala systemom komputerowym na „uczenie się” z danych bez jawnego programowania, identyfikując ukryte wzorce i korelacje. Przedsiębiorstwa wykorzystują szeroką gamę modeli ML do analizy danych o klientach, które mogą obejmować informacje demograficzne, historię zakupów, interakcje z marką na różnych platformach, a nawet dane nieustrukturyzowane, takie jak opinie w mediach społecznościowych czy skargi (GhorbanTanhaei et al., 2024).

W praktyce marketingowej stosuje się różnorodne algorytmy, z których każdy ma swoje specyficzne zastosowania, mocne strony i ograniczenia. Do najczęściej wykorzystywanych należą: regresja logistyczna, drzewa decyzyjne, lasy losowe, maszyny wektorów nośnych (Support Vector Machines - SVM) oraz sztuczne sieci neuronowe (Artificial Neural Networks - ANN), w tym zaawansowane architektury głębokiego uczenia, takie jak sieci Long Short-Term Memory (LSTM), szczególnie skuteczne w analizie danych sekwencyjnych i szeregów czasowych. Wybór odpowiedniego modelu jest decyzją strategiczną, uzależnioną

od charakteru problemu (np. klasyfikacja w predykcji churnu vs. regresja w prognozowaniu CLV), dostępności i jakości danych oraz pożądanego poziomu dokładności i interpretowalności wyników (Alhasan, 2025; Douaioui et al., 2024; William, 2025).

Rozwój i wdrażanie tych zaawansowanych modeli ujawniają kluczowe napięcie, z którym muszą mierzyć się strategicy marketingowi: kompromis między mocą predykcyjną a przejrzystością modelu. Prostsze modele, takie jak regresja liniowa czy drzewa decyzyjne, są stosunkowo łatwe do zinterpretowania, co pozwala menedżerom zrozumieć, dlaczego model wygenerował daną prognozę. Jednakże, często mają one trudności z uchwyceniem złożonych, nieliniowych zależności w danych. Z drugiej strony, zaawansowane modele, takie jak głębokie sieci neuronowe, wykazują znacznie wyższą skuteczność predykcyjną, ale ich wewnętrzne działanie jest często nieprzejrzyste, co określa się mianem „czarnej skrzynki” (William, 2025). Ta nieprzejrzystość stanowi poważne wyzwanie strategiczne. Marketingowiec potrzebuje nie tylko informacji, że dany klient zrezygnuje z usług, ale także zrozumienia przyczyn tego ryzyka, aby móc zaprojektować skuteczną interwencję. Poleganie na nieprzejrzystych, choć dokładnych modelach, może prowadzić do podejmowania „poprawnych” decyzji z „niewłaściwych” powodów, co rodzi ryzyko strategiczne i etyczne (Manzoor et al., 2024).

Ewolucja od prostych modeli statystycznych do złożonych sieci neuronowych pociąga za sobą również ewolucję w zakresie wymagań dotyczących danych. Skuteczne planowanie predykcyjne jest uzależnione od zdolności firmy do budowy zaawansowanej infrastruktury danych, zdolnej do integracji różnorodnych, wieloźródłowych i często strumieniowych danych w czasie rzeczywistym. Gotowość technologiczna przedsiębiorstw, zwłaszcza tych działających w SSE, jest zatem determinowana nie tylko przez ich zdolności algorytmiczne, ale przede wszystkim przez dojrzałość ich ekosystemu danych (William, 2025; Toorajipour et al., 2024) (tab. 7).

Tabela 7. Porównanie modeli uczenia maszynowego w marketingu predykcyjnym

Model	Główne zastosowanie	Mocne strony	Ograniczenia
Regresja logistyczna	Klasyfikacja binarna (np. predykcja churnu, scoring leadów, predykcja odpowiedzi na kampanię).	Prostota, szybkość, wysoka interpretowalność, niska złożoność obliczeniowa.	Niska skuteczność przy złożonych, nieliniowych wzorcach.

Model	Główne zastosowanie	Mocne strony	Ograniczenia
Drzewa decyzyjne	Segmentacja klientów, klasyfikacja, predykcja zachowań.	Intuicyjność, łatwość wizualizacji i interpretacji reguł decyzyjnych, zdolność do obsługi danych nieliniowych.	Podatność na przeuczenie, niestabilność (małe zmiany w danych mogą prowadzić do dużych zmian w strukturze drzewa).
Lasy losowe (Random Forest)	Predykcja churnu, prognozowanie CLV, zaawansowana segmentacja, scoring leadów.	Wysoka dokładność, odporność na przeuczenie dzięki agregacji wielu drzew, zdolność do obsługi dużej liczby zmiennych.	Mniejsza interpretowalność w porównaniu do pojedynczego drzewa decyzyjnego, większa złożoność obliczeniowa.
Maszyny wektorów nośnych (SVM)	Klasyfikacja (np. segmentacja klientów o wysokiej/niskiej wartości), regresja.	Wysoka skuteczność w przestrzeniach wielowymiarowych, efektywność w przypadkach, gdy liczba wymiarów jest większa niż liczba próbek.	Trudności w interpretacji, wrażliwość na wybór parametrów (jądra), duża złożoność obliczeniowa przy dużych zbiorach danych.
Sztuczne sieci neuronowe (ANN)	Prognozowanie popytu, predykcja CLV, systemy rekomendacyjne, rozpoznawanie wzorców w zachowaniach klientów.	Zdolność do modelowania bardzo złożonych, nieliniowych zależności, wysoka dokładność predykcyjna.	Efekt „czarnej skrzynki” (niska interpretowalność), wymagają dużych zbiorów danych do treningu, wysokie koszty obliczeniowe.
Sieci Long Short-Term Memory (LSTM)	Prognozowanie popytu w czasie, analiza sekwencji zachowań klientów (customer journey), predykcja trendów.	Wyjątkowa skuteczność w analizie danych sekwencyjnych i szeregów czasowych, zdolność do przechwytywania długoterminowych zależności.	Bardzo wysoka złożoność obliczeniowa i trudność w interpretacji, wymagają specjalistycznej wiedzy do implementacji i strojenia.

Źródło: Opracowanie własne na podstawie: Adwani, A. (2025). Predictive analytics for business strategy: Leveraging machine learning for competitive advantage. SSRN. https://www.researchgate.net/publication/393788373_Predictive_Analytics_for_Business_Strategy_Leveraging_Machine_Learning_for_Competitive_Advantage; Manzoor, A., Qureshi, M. A., Kidney, E., & Longo, L. (2024). A review on machine learning methods for customer churn prediction and recommendations for business practitioners. IEEE Access, 12, 70434-70463

Zastosowanie analityki predykcyjnej w marketingu nie ogranicza się do wykorzystania pojedynczych narzędzi czy algorytmów, lecz obejmuje budowę kompleksowego, zintegrowanego systemu modeli analitycznych, które wspierają zarządzanie całym cyklem życia klienta – od momentu jego pozyskania, przez rozwój i utrzymanie relacji, aż po proces retencji i ponownej aktywizacji. Tego rodzaju systemy integrują dane pochodzące z różnych źródeł, takich jak media społecznościowe, historia zakupów, interakcje z witryną internetową czy dane demograficzne, przekształcając je w użyteczne informacje biznesowe. Modele predykcyjne pozwalają między innymi przewidywać prawdopodobieństwo zakupu, określać ryzyko odejścia klienta, segmentować odbiorców według ich wartości i zachowań (Customer Lifetime Value, RFM analysis), a także prognozować skuteczność kampanii reklamowych czy optymalizować strategie cenowe. W efekcie przedsiębiorstwa mogą podejmować bardziej trafne, oparte na danych decyzje, minimalizować ryzyko i zwiększać zwrot z inwestycji w działania marketingowe. Co istotne, analityka predykcyjna pełni nie tylko funkcję operacyjną, lecz także strategiczną – umożliwia bowiem tworzenie długofalowych scenariuszy rozwoju relacji z klientami oraz identyfikację nowych możliwości rynkowych. Dzięki temu marketing przestaje być postrzegany jako koszt, a staje się kluczowym elementem budowy przewagi konkurencyjnej, opartej na wiedzy, personalizacji i zdolności do przewidywania przyszłych trendów konsumenckich.

Wartość życiowa klienta jest kluczowym wskaźnikiem, który pozwala ocenić długoterminową rentowność relacji z klientem i optymalizować strategie retencyjne. Tradycyjne metody, takie jak analiza RFM (Recency, Frequency, Monetary), choć użyteczne, mają charakter statyczny. Sztuczna inteligencja i uczenie maszynowe rewolucjonizują prognozowanie CLV, wprowadzając modele dynamiczne, które dostosowują się w czasie rzeczywistym do interakcji i zmian w zachowaniu klienta. Algorytmy takie jak regresja, drzewa decyzyjne czy sieci neuronowe pozwalają na znacznie dokładniejsze przewidywanie przyszłych zachowań zakupowych, co umożliwia firmom efektywniejszą alokację zasobów marketingowych, koncentrując się na klientach o najwyższym potencjale wartości (Akter et al., 2025).

Utrzymanie istniejącego klienta jest z reguły znacznie tańsze niż pozyskanie nowego, co czyni predykcję odejść jednym z najważniejszych zastosowań analityki predykcyjnej. Modele uczenia maszynowego, analizując dane demograficzne, transakcyjne i behawioralne, potrafią z dużym wyprzedzeniem zidentyfikować klientów zagrożonych rezygnacją z usług. Działają one jak system wczesnego ostrzegania, umożliwiając działom marketingu podejmowanie proaktywnych

działań retencyjnych, takich jak spersonalizowane oferty, ulepszona obsługa klienta czy dedykowane programy lojalnościowe. Skuteczność tych modeli, zwłaszcza w branżach o niskim progu zmiany dostawcy (np. telekomunikacja, e-commerce), ma bezpośredni wpływ na stabilność bazy klientów i rentowność przedsiębiorstwa (Manzoor et al., 2024).

Analityka predykcyjna pozwala na przejście od tradycyjnej, statycznej segmentacji opartej na danych demograficznych do dynamicznej mikro-segmentacji, która uwzględnia zachowania, preferencje i kontekst interakcji klienta z marką. Algorytmy ML potrafią automatycznie grupować klientów w oparciu o złożone wzorce, takie jak częstotliwość zakupów, zaangażowanie w różnych kanałach czy prawdopodobieństwo reakcji na określone typy kampanii. Taka precyzyjna segmentacja jest fundamentem skutecznej personalizacji, umożliwiając dostarczanie odpowiednich komunikatów właściwym grupom klientów we właściwym czasie (Kliś, 2025).

Modele predykcyjne odgrywają kluczową rolę w prognozowaniu popytu, co pozwala na optymalizację zarządzania zapasami, planowania promocji i strategii cenowych. Algorytmy ML, w szczególności te z rodziny deep learning, takie jak LSTM, są w stanie analizować złożone szeregi czasowe danych sprzedażowych i uwzględniać czynniki zewnętrzne, takie jak sezonowość, trendy rynkowe, działania konkurencji czy nawet warunki pogodowe, w celu generowania znacznie dokładniejszych prognoz niż metody tradycyjne. Prognozy te stanowią podstawę dla dynamicznej optymalizacji cen, gdzie AI w czasie rzeczywistym dostosowuje ceny produktów i usług w odpowiedzi na zmieniające się warunki rynkowe, maksymalizując przychody i rentowność (Kong et al., 2024; Basal et al., 2024).

Zastosowanie tych modeli w praktyce pokazuje, że ich prawdziwa siła nie leży w izolowanym działaniu, ale w ich wzajemnej integracji, która tworzy spójny system zarządzania wartością klienta. Proces ten można zobrazować następująco: zaawansowana segmentacja dzieli bazę klientów na precyzyjne grupy. Następnie, dla każdej z tych grup, modele prognozowania CLV określają ich długoterminową wartość, podczas gdy modele predykcji churnu identyfikują jednostki o podwyższonym ryzyku odejścia. Ta wiedza pozwala marketingowcom na priorytetyzację działań retencyjnych, decydując, ile zainwestować w utrzymanie danego klienta w oparciu o jego przewidywaną wartość. Oferty retencyjne mogą być następnie wyceniane przy użyciu modeli dynamicznego pricingu, które uwzględniają prognozy popytu na daną ofertę w konkretnym segmencie. W ten sposób planowanie marketingowe przekształca się z serii odrębnych

działań w zintegrowany, ciągle optymalizowany system, którego celem jest maksymalizacja wartości całego portfela klientów (Vargas et al., 2023; Liu et al., 2019).

Połączenie analityki predykcyjnej z personalizacją stanowi rdzeń nowoczesnej strategii marketingowej. Predykcja dostarcza informacji o tym, co klient prawdopodobnie zrobi, zechce lub kupi, podczas gdy personalizacja wykorzystuje te informacje do stworzenia zindywidualizowanego doświadczenia, które na tę prognozę odpowiada. Mechanizm ten stanowi ewolucję od prostych, opartych na regułach systemów („jeśli klient obejrzał produkt X, wyślij mu e-mail z ofertą Y”) do zaawansowanych, adaptacyjnych systemów napędzanych przez uczenie maszynowe. Nowoczesne platformy personalizacyjne analizują w czasie rzeczywistym ogromne ilości danych – historię przeglądania, wzorce zakupowe, interakcje w mediach społecznościowych, a nawet kontekst sytuacyjny (np. pora dnia, lokalizacja) – aby dynamicznie dostosowywać treści, rekomendacje produktowe i oferty (Dorgbefu, 2021).

Skuteczność tego podejścia została potwierdzona w licznych badaniach i wdrożeniach. Studium przypadku detalisty z branży modowej wykazało, że zintegrowany system, łączący silnik rekomendacyjny Amazon Personalize z dostrojonymi modelami językowymi GPT do dynamicznego generowania treści, przyniósł wzrost konwersji z działań upsellingowych o 27% oraz poprawę postrzeganej jakości personalizacji o 18%. Inne badania wskazują, że niszowi detaliści, stosując narzędzia AI, odnotowali wzrost wskaźnika retencji klientów o 30%. Giganci technologiczni, tacy jak Amazon, Netflix czy Spotify, ustanowili nowe standardy w branży, wykorzystując personalizację do zwiększania zaangażowania i lojalności klientów. Personalizacja nie ogranicza się już tylko do rekomendacji produktowych; obejmuje dynamiczne dostosowywanie całych stron internetowych, treści e-maili marketingowych i kampanii reklamowych do indywidualnych profili użytkowników (Beem, 2025; Bell et al., 2024).

Najnowsze trendy wskazują na konwergencję analityki predykcyjnej z generatywną sztuczną inteligencją. Wspomniane studium przypadku detalisty modowego jest tego doskonałym przykładem. Firma wykorzystwała predykcyjny silnik rekomendacyjny (Amazon Personalize) do określenia, jakie produkty zarekomendować, a następnie użyła modeli generatywnych (GPT) do stworzenia spersonalizowanego opisu marketingowego dla tej rekomendacji. To zaawansowane zastosowanie pokazuje, że przyszłość personalizacji leży nie tylko w trafności rekomendacji, ale także w adekwatności komunikacji. Planowanie marketingowe musi zatem uwzględniać strategię treści na poziomie algorytmicznym, trenując modele generatywne na podstawie głosu marki, najsukuteczniejszych tekstów

reklamowych i sentymentu klientów, aby dostarczać personalizację, która jest nie tylko dokładna, ale również przekonująca i spójna z wizerunkiem marki (Beem, 2025).

Pomimo ogromnego potencjału, wdrażanie analityki predykcyjnej w marketingu wiąże się z szeregiem wyzwań, zarówno technicznych, jak i etycznych. Do barier technicznych należą przede wszystkim problemy z jakością i dostępnością danych, wysokie koszty implementacji i utrzymania zaawansowanych systemów oraz niedobór wykwalifikowanych specjalistów, którzy potrafiliby skutecznie zarządzać tymi technologiami. Jednak to dylematy etyczne stanowią największe i najtrudniejsze do rozwiązania wyzwanie, które może zaważyć na długoterminowym sukcesie strategii opartych na danych.

Jednym z najpoważniejszych problemów jest stronniczość algorytmiczna. Analiza przypadku kampanii reklamowej promującej zatrudnienie w branżach STEM (nauka, technologia, inżynieria, matematyka) wykazała, że algorytm, którego celem była optymalizacja kosztów, wyświetlał reklamę znacznie częściej mężczyznom niż kobietom. Działo się tak, ponieważ młode kobiety są cenną i droższą grupą docelową w reklamie cyfrowej. Algorytm, dążąc do maksymalizacji efektywności kosztowej, nieumyślnie wygenerował dyskryminujący rezultat, mimo że intencją kampanii była neutralność płciowa. Ten przykład ilustruje fundamentalną prawdę: błędy etyczne często nie są wynikiem złych intencji, ale niezamierzoną konsekwencją wąsko zdefiniowanych celów biznesowych. Stronniczość nie jest „błędem” w kodzie, który można naprawić, ale emergentną właściwością systemu optymalizującego pojedynczy wskaźnik (np. koszt, konwersja) bez uwzględnienia ograniczeń etycznych. To z kolei implikuje, że proces planowania strategicznego musi ewoluować. Marketingowcy nie mogą po prostu przekazać celów biznesowych analitykom danych; muszą aktywnie uczestniczyć w definiowaniu „ograniczeń sprawiedliwości” i etycznych ram dla działania algorytmów. Kluczowe pytanie strategiczne brzmi nie tylko „jaki jest nasz cel ROI?”, ale „jaki jest nasz cel ROI, przy założeniu, że nasze kampanie dotrą do reprezentatywnej grupy odbiorców?” (Lambrecht & Tucker, 2018; Akter et al., 2023).

Inne istotne kwestie etyczne to brak przejrzystości (problem „czarnej skrzynki”), ochrona prywatności danych i budowanie zaufania konsumentów. Konsumenti mają prawo wiedzieć, w jaki sposób ich dane są wykorzystywane i na jakiej podstawie otrzymują spersonalizowane komunikaty. W dobie rosnącej świadomości społecznej i regulacji prawnych, takich jak RODO (GDPR), etyczne zarządzanie danymi przestaje być jedynie kwestią zgodności z prawem, a staje się kluczowym elementem wartości marki i źródłem przewagi konkurencyjnej.

Firma, która w sposób transparentny komunikuje, jak wykorzystuje dane do tworzenia wartości dla klienta – bez uciekania się do manipulacji – buduje silniejsze i bardziej lojalne relacje. Dlatego planowanie marketingowe musi traktować etykę nie jako ograniczenie dla rentowności, ale jako inwestycję w długoterminowy, niematerialny kapitał, jakim jest zaufanie klientów (Oluwafemi et al., 2021).

Integracja analityki predykcyjnej i personalizacji przestała być niszową zdolnością technologiczną, a stała się kluczowym imperatywem strategicznym każdego nowoczesnego przedsiębiorstwa, które dąży do osiągnięcia oraz utrzymania przewagi konkurencyjnej w warunkach gospodarki opartej na danych. Zdolność do wykorzystania analityki predykcyjnej w procesie podejmowania decyzji stanowi obecnie jedno z najważniejszych źródeł wzrostu efektywności organizacyjnej i innowacyjności. Dzięki zaawansowanym technikom analizy danych przedsiębiorstwa mogą nie tylko przewidywać przyszłe zachowania klientów, lecz także projektować spersonalizowane doświadczenia, które odpowiadają ich indywidualnym potrzebom, preferencjom i kontekstowi interakcji z marką. Współczesny marketing coraz częściej opiera się na koncepcji proaktywnego zarządzania relacjami z klientem, w której dane historyczne i behawioralne są wykorzystywane do formułowania predykcji w czasie rzeczywistym. W praktyce oznacza to, że przedsiębiorstwa potrafią identyfikować potencjalne potrzeby konsumentów jeszcze zanim zostaną one przez nich wyrażone, co umożliwia tworzenie dynamicznie dopasowanej oferty. Tego rodzaju podejście nie tylko zwiększa skuteczność działań sprzedażowych i marketingowych, ale również wzmacnia satysfakcję klientów oraz ich lojalność wobec marki. Zdolność do przewidywania zachowań konsumentów i dostarczania im spersonalizowanych, wartościowych doświadczeń staje się zatem strategicznym zasobem, który łączy wymiar technologiczny z relacyjnym. Firmy, które skutecznie integrują analitykę predykcyjną z systemami personalizacji, przekształcają dane w wiedzę, a wiedzę – w decyzje o wysokiej precyzji, przyczyniające się do wzrostu wartości przedsiębiorstwa. Organizacje o wysokim stopniu dojrzałości analitycznej osiągają wyższy poziom efektywności operacyjnej, rentowności i elastyczności strategicznej, co w długiej perspektywie prowadzi do uzyskania trwałej przewagi konkurencyjnej. Warto podkreślić, że rola analityki predykcyjnej nie ogranicza się jedynie do sfery operacyjnej marketingu. Coraz częściej stanowi ona fundament kultury organizacyjnej opartej na danych, w której decyzje strategiczne – od projektowania nowych produktów po zarządzanie doświadczeniem klienta – są podejmowane na podstawie analitycznych wniosków, a nie intuicji. W ten sposób integracja analityki i personalizacji wspiera procesy zrównoważonego wzrostu przedsiębiorstw, umożliwiając im nie

tylko skuteczne reagowanie na zmiany rynkowe, lecz także ich przewidywanie i kształtowanie. Podsumowując, synergiczne połączenie analityki predykcyjnej i personalizacji staje się jednym z kluczowych czynników sukcesu organizacji w erze transformacji cyfrowej. Firmy, które potrafią zintegrować te dwa elementy w spójną strategię zarządzania relacjami z klientem, uzyskują zdolność do tworzenia unikalnej wartości, zwiększania zaufania i lojalności konsumentów, a tym samym – budowania trwałych podstaw konkurencyjności w długim okresie (Adesina et al., 2024).

Jednakże, aby w pełni wykorzystać potencjał technologii opartych na sztucznej inteligencji i analityce predykcyjnej, przedsiębiorstwa muszą rozwijać kompetencje w trzech wzajemnie powiązanych i komplementarnych obszarach. Skuteczność wdrażania rozwiązań AI w marketingu i zarządzaniu zależy nie tylko od samej dostępności technologii, lecz przede wszystkim od poziomu dojrzałości organizacyjnej oraz zdolności do zintegrowania wiedzy technologicznej, strategicznej i etycznej w ramach spójnego modelu działania. Po pierwsze, kluczowym filarem jest gotowość technologiczna, obejmująca zarówno infrastrukturę technologiczną, jak i zdolność organizacji do jej efektywnego wykorzystania. Nie sprowadza się ona wyłącznie do wdrożenia zaawansowanych algorytmów uczenia maszynowego czy systemów rekomendacyjnych, lecz wymaga stworzenia solidnych fundamentów w postaci infrastruktury danych – zdolnej do gromadzenia, integracji, przetwarzania i analizy informacji z wielu źródeł w czasie rzeczywistym. Tylko w ten sposób możliwe jest budowanie systemów predykcyjnych, które dostarczają wiarygodnych i użytecznych wniosków wspierających proces decyzyjny. Gotowość technologiczna wiąże się zatem bezpośrednio z rozwojem kompetencji cyfrowych pracowników, kulturą organizacyjną sprzyjającą innowacjom oraz umiejętnością adaptacji do zmieniających się technologii. Po drugie, niezbędny jest rozwinięty zmysł strategiczny, rozumiany jako zdolność do przekładania wyników analiz predykcyjnych na realne, kreatywne i angażujące działania marketingowe. Obejmuje on umiejętność interpretacji danych w kontekście szerszych celów biznesowych oraz przekształcania prognoz w strategię komunikacyjne, które skutecznie rezonują z wartościami i oczekiwaniami klientów. W tym wymiarze kluczowe znaczenie ma integracja analizy danych z procesami planowania marketingowego, zarządzania marką oraz projektowania doświadczeń klienta. Firmy, które potrafią połączyć analitykę z wizją strategiczną, przekształcają dane w przewagę rynkową, a wiedzę – w kapitał relacyjny. Po trzecie, równie istotne pozostaje zarządzanie etyczne, które stanowi nieodzowny element odpowiedzialnego wykorzystywania technologii predykcyjnych. Wraz z rosnącą zdolnością

systemów AI do analizowania zachowań i preferencji konsumentów pojawiają się nowe wyzwania związane z ochroną prywatności, przejrzystością algorytmów oraz zapobieganiem dyskryminacji danych. Etyczne zarządzanie oznacza zatem świadome i odpowiedzialne korzystanie z technologii w sposób, który respektuje prawa jednostki, zapewnia sprawiedliwość decyzyjną oraz buduje zaufanie między firmą a klientem. Równowaga między innowacyjnością a etyką staje się jednym z kluczowych czynników warunkujących długofalową akceptację społeczną i trwałość przewagi konkurencyjnej przedsiębiorstw. Pełne wykorzystanie potencjału sztucznej inteligencji w marketingu wymaga holistycznego podejścia, łączącego zaawansowaną infrastrukturę technologiczną, strategiczne myślenie oraz odpowiedzialne zarządzanie etyczne. Dopiero współistnienie tych trzech komponentów pozwala organizacjom na zbudowanie zrównoważonego modelu przewagi konkurencyjnej opartej na danych, zaufaniu i innowacyjności (Oluwafemi et al., 2021).

Przyszłość planowania marketingowego będzie należeć do organizacji, które potrafią harmonijnie połączyć trzy kluczowe filary: gotowość technologiczną, zmysł strategiczny oraz zarządzanie etyczne. Sukces w erze sztucznej inteligencji nie będzie wynikał jedynie z posiadania najnowocześniejszych narzędzi analitycznych, lecz z umiejętności ich zintegrowania z ludzką kreatywnością, odpowiedzialnością społeczną i zdolnością adaptacyjną organizacji. Przyszły model marketingu będzie oparty na synergii między technologią a człowiekiem, w której automatyzacja, personalizacja i analityka predykcyjna współtworzą spójny ekosystem decyzyjny wspierający innowacyjność oraz zrównoważony rozwój. Wymaga to jednak nie tylko wdrożenia nowych technologii, lecz także przeformułowania tradycyjnych struktur organizacyjnych i ról zawodowych. Wzrasta zapotrzebowanie na tzw. „inżynierów marketingu” – specjalistów łączących kompetencje z zakresu analityki danych, programowania, uczenia maszynowego oraz klasycznego marketingu strategicznego. Ich zadaniem będzie nie tylko obsługa narzędzi technologicznych, lecz również tłumaczenie wyników analiz na język biznesu i strategii rynkowej. Tacy eksperci staną się kluczowym ogniwem między światem danych a światem decyzji, pomagając organizacjom skutecznie wykorzystywać potencjał predykcji i automatyzacji do tworzenia wartości dla klienta. W konsekwencji proces planowania strategicznego ulegnie transformacji w kierunku interdyscyplinarnej współpracy, w której analitycy danych, specjaliści ds. marketingu, etycy i liderzy biznesowi będą wspólnie kształtować decyzje oparte na danych. Takie zespoły, działające w modelu cross-funkcyjnym, będą nie tylko optymalizować skuteczność kampanii, lecz także dbać o zachowanie równowagi

między efektywnością technologiczną a odpowiedzialnością społeczną. Jednym z kluczowych wyzwań nadchodzącej dekady będzie zarządzanie etycznymi implikacjami wykorzystania danych – w tym prywatnością, transparentnością i sprawiedliwością algorytmiczną. W praktyce oznacza to, że strategiczny marketing przyszłości stanie się dziedziną z pogranicza technologii, psychologii, ekonomii behawioralnej i etyki, w której decyzje nie mogą być oparte wyłącznie na prognozach algorytmicznych, lecz muszą uwzględniać czynniki społeczne, kulturowe i środowiskowe. Organizacje, które będą w stanie połączyć moc predykcji z empatią i odpowiedzialnością, zyskają trwałą przewagę w budowaniu zaufania i lojalności klientów. Dla przedsiębiorstw działających w Specjalnych Strefach Ekonomicznych, rozwój tak zintegrowanej zdolności stanie się kluczowym wyznacznikiem gotowości do konkutowania na globalnym rynku. Strefy te, tradycyjnie koncentrujące się na przyciąganiu inwestycji dzięki zachętom fiskalnym, będą musiały ewoluować w kierunku inteligentnych ekosystemów innowacji, wspierających rozwój kompetencji cyfrowych, współpracę międzysektorową i etyczne wdrażanie technologii AI w praktyce biznesowej. W tym ujęciu, konkurencyjność firm zlokalizowanych w SSE będzie coraz bardziej zależeć od ich zdolności do tworzenia interdyscyplinarnych zespołów, integrujących wiedzę technologiczną z rozumieniem zachowań konsumentów oraz zasadami odpowiedzialnego zarządzania. Podsumowując, marketing przyszłości będzie oparty na współdziałaniu ludzi i maszyn, w którym dane staną się wspólnym językiem łączącym analitykę, strategię i etykę. Firmy, które zrozumieją tę nową logikę organizacyjną i nauczą się efektywnie łączyć kompetencje technologiczne z humanistycznym podejściem do klienta, nie tylko utrzymają konkurencyjność, ale również staną się liderami zmian w erze gospodarki cyfrowej (Hunter, 2025).

Implementacja działań marketingowych i przewaga konkurencyjna przedsiębiorstw

W erze wszechobecnej transformacji cyfrowej, która fundamentalnie redefiniuje krajobraz gospodarczy, sztuczna inteligencja ewoluuje z roli narzędzia wspierającego do strategicznego katalizatora, który umożliwia przedsiębiorstwom nie tylko optymalizację bieżących operacji, ale przede wszystkim budowanie i utrzymywanie trwałej przewagi konkurencyjnej. Wdrożenie AI w strategii marketingowej przestało być innowacyjną opcją, a stało się imperatywem dla organizacji aspirujących do roli liderów rynkowych. Zdolność do wykorzystania potencjału AI determinuje zdolność firmy do adaptacji, innowacji i, w konsekwencji,

do przetrwania na coraz bardziej dynamicznym i wymagającym rynku (Dzreke, 2025; Gonçalves et al., 2022).

Tradycyjne ujęcie przewagi konkurencyjnej, definiowanej jako ogół czynników wyróżniających przedsiębiorstwo i zwiększających atrakcyjność jego oferty w oczach klientów w stosunku do konkurentów, ulega obecnie głębokiej rekonfiguracji pod wpływem postępu technologicznego. Współczesne podejście do konkurencyjności odchodzi od statycznego postrzegania przewagi, skupionego na posiadaniu zasobów trudno imitowalnych, w kierunku koncepcji dynamicznej przewagi konkurencyjnej, która wynika z umiejętności szybkiej adaptacji, innowacyjności oraz zdolności do wykorzystania wiedzy i technologii w zmieniającym się otoczeniu rynkowym. Przewaga konkurencyjna we współczesnej gospodarce może mieć charakter schumpeterowski – opierać się na posiadaniu bardziej zaawansowanej, efektywniejszej technologii, umożliwiającej tworzenie nowych modeli biznesowych, produktów i usług. Tego rodzaju przewaga wynika z procesów twórczej destrukcji, w których innowacje technologiczne wypierają dotychczasowe rozwiązania, prowadząc do przekształceń struktur rynkowych. Jednocześnie autorzy podkreślają znaczenie zdolności organizacji do osiągania relatywnie szybszego wzrostu produktywności w porównaniu z konkurentami – co odzwierciedla kluczową rolę wiedzy, kompetencji cyfrowych i elastyczności w zarządzaniu zasobami. W tym kontekście sztuczna inteligencja jawi się jako technologia o charakterze ogólnego zastosowania, której wpływ wykracza poza pojedyncze sektory, oddziałując systemowo na całą gospodarkę. AI umożliwia automatyzację procesów, przyspieszenie analizy danych, zwiększenie efektywności operacyjnej oraz personalizację oferty – co w połączeniu tworzy nową jakość przewagi konkurencyjnej zarówno na poziomie mikro-, jak i makroekonomicznym. Zastosowanie tej technologii pozwala przedsiębiorstwom nie tylko osiągać wyższy poziom produktywności, ale również budować zdolność do ciągłego uczenia się i innowacji, co stanowi fundament długookresowej konkurencyjności. Transformacja cyfrowa, której sztuczna inteligencja stanowi rdzeń, ma charakter wielowymiarowy – łączy aspekty technologiczne, organizacyjne, społeczne, ekonomiczne i środowiskowe. Jej oddziaływanie obejmuje zarówno mikrostruktury przedsiębiorstw, jak i makroukłady gospodarcze, determinując pozycję konkurencyjną krajów i regionów. W ujęciu makroekonomicznym technologie cyfrowe przyczyniają się do wzrostu efektywności alokacji zasobów, poprawy produktywności pracy oraz rozwoju nowych sektorów gospodarki opartej na danych. W wymiarze mikroekonomicznym umożliwiają firmom budowanie elastycznych modeli działania, zwiększanie wartości dodanej oraz tworzenie zindywidualizowanych

doświadczeń klienta, co przekłada się na trwałą i trudną do skopiowania przewagę konkurencyjną. W rezultacie, w erze gospodarki cyfrowej, przewaga konkurencyjna coraz częściej opiera się nie na posiadaniu materialnych zasobów, lecz na zdolności do efektywnego wykorzystania technologii, wiedzy i danych. Organizacje, które potrafią zintegrować sztuczną inteligencję z procesami strategicznymi, stają się liderami zmian, zdolnymi do kształtowania nowych standardów efektywności i innowacyjności w skali całych sektorów (Kowalski et al., 2023).

Fundamentem, na którym przedsiębiorstwa mogą budować zaawansowane zdolności oparte na sztucznej inteligencji, jest odpowiedni poziom gotowości technologicznej. Współczesna gospodarka, charakteryzująca się wysokim stopniem digitalizacji i globalnej współzależności, wymaga od organizacji nie tylko dostępu do nowoczesnych technologii, lecz także zdolności do ich skutecznej implementacji i integracji z procesami biznesowymi. Technologia stanowi dziś kluczowy czynnik umożliwiający przedsiębiorstwom nie tylko konkutowanie, ale i utrzymanie długookresowej zdolności do adaptacji w dynamicznie zmieniającym się otoczeniu rynkowym. Gotowość technologiczna odzwierciedla więc poziom zaawansowania infrastruktury cyfrowej, kompetencji technicznych, kultury organizacyjnej sprzyjającej innowacjom oraz zdolności do wykorzystania danych i automatyzacji w procesach decyzyjnych. W zglobalizowanym środowisku gospodarczym, w którym przewagi konkurencyjne coraz częściej wynikają z umiejętności przetwarzania informacji i wdrażania inteligentnych rozwiązań, gotowość technologiczna staje się jednym z filarów konkurencyjności zarówno przedsiębiorstw, jak i całych gospodarek. Krajowe i regionalne różnice w tym zakresie wpływają bezpośrednio na zdolność przyciągania inwestycji, rozwój innowacji oraz tempo transformacji cyfrowej. Wysoki poziom gotowości technologicznej sprzyja szybszemu wdrażaniu rozwiązań opartych na sztucznej inteligencji, automatyzacji procesów produkcyjnych oraz budowie ekosystemów innowacji, które generują efekt synergii między sektorem prywatnym, publicznym i naukowym. W kontekście Specjalnych Stref Ekonomicznych w Polsce, problematyka ta nabiera szczególnego znaczenia. Tradycyjnie strefy te przyciągały inwestorów głównie poprzez instrumenty finansowe, takie jak ulgi podatkowe czy niższe koszty operacyjne, pełniąc funkcję katalizatora rozwoju regionalnego. Jednak w warunkach gospodarki opartej na wiedzy i technologii, czynniki kosztowe przestają być wystarczającym bodźcem do utrzymania konkurencyjności. O długoterminowej atrakcyjności i zdolności SSE do generowania wartości dodanej coraz częściej decydują elementy związane z innowacyjnością, infrastrukturą badawczo-rozwojową, dostępem do wykwalifikowanej kadry oraz możliwością

współpracy z ośrodkami naukowymi. Współczesne SSE powinny zatem ewoluować w kierunku tzw. inteligentnych stref przemysłowych, które łączą tradycyjne instrumenty wsparcia inwestycji z nowoczesnymi technologiami cyfrowymi i rozwiązaniami z zakresu Przemysłu 4.0. W takich ekosystemach rośnie znaczenie infrastruktury teleinformatycznej, centrów danych, laboratoriów innowacji oraz programów wspierających rozwój kompetencji cyfrowych pracowników. Gotowość technologiczna staje się więc nie tylko miernikiem dojrzałości organizacyjnej przedsiębiorstw funkcjonujących w strefach, ale także kluczowym czynnikiem determinującym ich zdolność do trwałego wzrostu produktywności i konkurencyjności. A zatem, w dobie przyspieszonej transformacji cyfrowej, gotowość technologiczna jest nieodzownym warunkiem rozwoju zdolności opartych na sztucznej inteligencji. Stanowi ona podstawę budowania przewagi konkurencyjnej opartej na wiedzy, innowacjach i efektywności procesów. W przypadku Specjalnych Stref Ekonomicznych, inwestycje w nowoczesną infrastrukturę technologiczną i kompetencje cyfrowe stają się nie tylko środkiem do przyciągania inwestorów, lecz także koniecznym elementem zapewnienia ich długofalowej atrakcyjności i odporności na zmiany globalnych trendów gospodarczych (Uren & Edwards, 2023; Michałek, 2025).

Obserwujemy fundamentalną ewolucję źródeł przewagi konkurencyjnej – od przewagi kosztowej do przewagi kognitywnej. Historycznie, wiele przedsiębiorstw, w tym te działające w SSE, budowało swoją pozycję na tradycyjnych, kosztowych czynnikach, które w perspektywie długoterminowej nie stanowią trwałej podstawy do konkurowania na rynkach międzynarodowych. Sztuczna inteligencja, definiowana jako zdolność systemu do prawidłowego interpretowania danych, uczenia się na ich podstawie oraz wykorzystywania tej wiedzy do osiągnięcia celów poprzez elastyczne dostosowanie, wprowadza nowy paradygmat. Przewaga kognitywna to zdolność organizacji do szybszego uczenia się, przewidywania i adaptacji niż konkurenci, oparta na ciągłej analizie danych w czasie rzeczywistym. Implementacja AI w marketingu nie jest zatem jedynie optymalizacją kosztów, lecz fundamentalną zmianą modelu biznesowego w kierunku organizacji „uczącej się”. Przewaga nie leży już tylko w tym, co firma posiada (aktywa, ulgi podatkowe), ale w tym, co wie i jak szybko potrafi tę wiedzę przekuć w skuteczne działania rynkowe (Climent et al., 2024; Musiał, 2019).

Sztuczna inteligencja rewolucjonizuje paradygmat zarządzania relacjami z klientem (CRM), umożliwiając fundamentalne przejście od tradycyjnej, statycznej segmentacji opartej na danych demograficznych do dynamicznej, behawioralnej hiperpersonalizacji realizowanej na masową skalę. AI umożliwia organizacjom

tworzenie systemów CRM nowej generacji, które uczą się na podstawie interakcji z klientami, rozpoznając nie tylko ich bieżące potrzeby, lecz także potencjalne motywacje i przyszłe zachowania zakupowe. Nowoczesne rozwiązania oparte na uczeniu maszynowym pozwalają na integrację różnorodnych źródeł danych – od aktywności online i historii transakcji, po dane z mediów społecznościowych i kontekstowe sygnały w czasie rzeczywistym. Dzięki temu możliwe jest konstruowanie szczegółowych profili klientów, które ewoluują w czasie i pozwalają markom dostosowywać komunikację, ofertę i doświadczenia w sposób dynamiczny i predykcyjny. Tego rodzaju podejście przekształca CRM z systemu zarządzania danymi w inteligentny ekosystem wspierający personalizację i automatyczne podejmowanie decyzji. Autorzy podkreślają, że hiperpersonalizacja nie ogranicza się jedynie do dopasowywania treści marketingowych, lecz obejmuje również indywidualne projektowanie całej ścieżki klienta. Sztuczna inteligencja umożliwia automatyczne modyfikowanie interfejsów, rekomendacji produktów, ofert cenowych czy sekwencji kontaktów w zależności od kontekstu użytkownika. Oznacza to, że każda interakcja staje się unikalnym doświadczeniem, dostosowanym do aktualnego nastroju, lokalizacji czy wcześniejszych preferencji klienta. Kluczową przewagą takiego podejścia jest zdolność do przewidywania zachowań konsumenckich z wysokim poziomem dokładności. Modele predykcyjne, oparte na analizie danych behawioralnych, umożliwiają identyfikowanie klientów o wysokim prawdopodobieństwie zakupu, rezygnacji lub zmiany preferencji, co pozwala firmom działać proaktywnie. W ten sposób marketing przestaje być reaktywny, a staje się predykcyjny i empatyczny, budując głębsze relacje z klientami na bazie zrozumienia ich indywidualnych potrzeb. Rozwój hiperpersonalizacji wymaga również nowego podejścia do etyki danych. Firmy powinny zapewniać przejrzystość procesów analitycznych i dbać o równowagę między personalizacją a prywatnością użytkownika. W przeciwnym razie ryzykują utratę zaufania, które jest kluczowym elementem współczesnych relacji konsumenckich. Zdolność do dostarczania unikalnych, dopasowanych w czasie rzeczywistym doświadczeń każdemu klientowi staje się fundamentem budowania lojalności i jednym z najważniejszych źródeł przewagi konkurencyjnej opartej na zróżnicowaniu. W erze marketingu predykcyjnego hiperpersonalizacja przestaje być luksusem technologicznym – staje się nowym standardem obsługi i relacyjnego podejścia do klienta, definiując przyszłość zarządzania doświadczeniem konsumenta (Ramya & Kumar, 2025).

Mechanizmy napędzające transformację współczesnego marketingu i zarządzania relacjami z klientem opierają się na zaawansowanych zdolnościach

analitycznych sztucznej inteligencji, które umożliwiają pogłębione rozumienie zachowań konsumentów, a także dynamiczne reagowanie na ich potrzeby w czasie rzeczywistym. Istota przewagi konkurencyjnej w nowoczesnym marketingu nie wynika już z samego dostępu do danych, lecz z umiejętności ich interpretacji, integracji i zastosowania w sposób predykcyjny. Po pierwsze, algorytmy uczenia maszynowego oraz techniki przetwarzania języka naturalnego (NLP) stanowią fundament współczesnej segmentacji klientów. AI potrafi analizować ogromne zbiory danych pochodzące z różnych źródeł – historii zakupów, aktywności w kanałach cyfrowych, ścieżek nawigacji internetowej czy interakcji w mediach społecznościowych – w celu identyfikowania nieoczywistych, ukrytych wzorców w zachowaniach konsumentów. Dzięki takim możliwościom powstają mikrosegmenty, a w najbardziej zaawansowanych przypadkach segmenty jednoosobowe, które dynamicznie ewoluują w czasie rzeczywistym. Tego rodzaju personalizacja predykcyjna wykracza poza możliwości tradycyjnych metod analizy danych, umożliwiając dopasowanie oferty do indywidualnych preferencji i kontekstu każdej jednostki. Po drugie, sztuczna inteligencja umożliwia dynamiczne generowanie i dostosowywanie treści marketingowych w oparciu o aktualne dane kontekstowe. Jednym z przykładów są rozwiązania typu Open Time Content, pozwalają na personalizację zawartości komunikatów w momencie ich otwarcia, a nie wysyłki. Dzięki temu każda wiadomość, reklama lub rekomendacja jest zawsze aktualna, kontekstowo dopasowana i maksymalnie trafna z perspektywy odbiorcy. Tego rodzaju adaptacyjność komunikacji prowadzi do wzrostu skuteczności działań marketingowych, ale również buduje wrażenie autentycznego dialogu z klientem, co sprzyja lojalności i długofalowej relacji z marką. Po trzecie, sercem hiperpersonalizacji pozostają inteligentne systemy rekomendacyjne, wykorzystujące algorytmy głębokiego uczenia oraz modele oparte na sieciach neuronowych do przewidywania przyszłych potrzeb użytkowników. Systemy te analizują w czasie rzeczywistym zachowania klientów, ich historię przeglądania, decyzje zakupowe oraz dane kontekstowe, generując propozycje produktów i usług, które nie tylko odpowiadają dotychczasowym preferencjom, lecz również antycypują potrzeby, których klient nie jest jeszcze w pełni świadomy. W ten sposób rekomendacje stają się nie tyle odzwierciedleniem przeszłości, co projekcją przyszłego potencjału zakupowego, co znacząco zwiększa skuteczność komunikacji i wartość relacji z klientem. Integracja tych trzech komponentów – analizy behawioralnej, dynamicznej personalizacji treści oraz predykcyjnych rekomendacji – prowadzi do powstania ekosystemu danych w czasie rzeczywistym, w którym decyzje marketingowe podejmowane są automatycznie, w oparciu o zmieniające się sygnały

z otoczenia. Taki model działania pozwala organizacjom na bieżące dostosowywanie strategii do aktualnych trendów i emocji konsumentów, eliminując opóźnienia informacyjne i zwiększając precyzję działań. Zaawansowane zdolności analityczne sztucznej inteligencji stanowią fundament współczesnej hiperpersonalizacji i redefiniują sposób, w jaki przedsiębiorstwa rozumieją, angażują i obsługują swoich klientów. AI przekształca marketing z funkcji reaktywnej w system dynamicznego uczenia się o konsumencie, w którym granica między analizą danych a doświadczeniem klienta ulega zatarciu. Organizacje, które potrafią skutecznie wykorzystać potencjał tych technologii, zyskują zdolność do tworzenia długofalowej przewagi opartej na wiedzy, elastyczności i głębokim zrozumieniu potrzeb odbiorcy (Ramya et al., 2025).

Bezpośrednią konsekwencją wdrożenia mechanizmów opartych na sztucznej inteligencji jest głęboka transformacja doświadczenia klienta, które coraz częściej staje się jednym z kluczowych obszarów budowania przewagi konkurencyjnej w gospodarce cyfrowej. Współcześni konsumenci oczekują od marek obsługi, która jest nie tylko szybka i dostępna, ale przede wszystkim spersonalizowana, kontekstowa i proaktywna. Technologie oparte na AI pozwalają organizacjom sprostać tym rosnącym wymaganiom, umożliwiając automatyczną analizę danych o klientach, rozpoznawanie wzorców zachowań oraz przewidywanie ich przyszłych potrzeb. Zarządzanie doświadczeniem klienta przestaje być jedynie elementem komunikacji marketingowej, a staje się centralnym filarem strategii biznesowej nowoczesnych przedsiębiorstw. AI umożliwia tworzenie modeli obsługi, które łączą w sobie elastyczność, indywidualne podejście oraz zdolność do reagowania w czasie rzeczywistym. W tym kontekście doświadczenie klienta nie jest już postrzegane jako efekt pojedynczej interakcji, lecz jako całościowy proces obejmujący wszystkie punkty styku z marką – od pierwszego kontaktu po obsługę posprzedażową. Jednym z najważniejszych obszarów zastosowania sztucznej inteligencji w CX jest rozwój inteligentnych chatbotów i wirtualnych asystentów, które dzięki technikom przetwarzania języka naturalnego potrafią rozumieć intencje, emocje i kontekst rozmowy. Jak podkreśla Tego typu rozwiązania znacząco poprawiają jakość komunikacji z klientem, umożliwiając szybkie reagowanie na zapytania, rozwiązywanie problemów i utrzymanie pozytywnego wizerunku marki. Wykorzystanie analizy sentymentu pozwala dodatkowo dostosować ton i sposób odpowiedzi do nastroju klienta, co wzmacnia wrażenie empatii i autentycznego zaangażowania ze strony firmy. Sztuczna inteligencja w obszarze CX wspiera także rozwój strategii omnichannel, umożliwiając integrację różnych kanałów komunikacji – od czatu i mediów społecznościowych

po kontakt telefoniczny czy e-mailowy – w jeden spójny system. Dzięki temu klienci doświadczają płynnych i konsekwentnych interakcji z marką, niezależnie od wybranego sposobu kontaktu. Tego rodzaju spójność komunikacyjna stanowi dziś kluczowy element budowania zaufania i lojalności klientów. Z perspektywy organizacyjnej wdrożenie rozwiązań AI w obszarze obsługi klienta przynosi podwójne korzyści. Z jednej strony pozwala zwiększyć efektywność operacyjną i odciążyć personel od rutynowych zadań, z drugiej – umożliwia pogłębienie relacji z klientem poprzez personalizację i szybsze reagowanie na jego potrzeby. Organizacje, które potrafią włączyć sztuczną inteligencję w strategiczne zarządzanie doświadczeniem klienta, osiągają wyższy poziom satysfakcji, lojalności i pozytywnego postrzegania marki. Sztuczna inteligencja redefiniuje sposób, w jaki przedsiębiorstwa projektują i zarządzają doświadczeniem klienta. Przestaje być jedynie narzędziem wspierającym obsługę, a staje się strategicznym komponentem relacji z rynkiem, łączącym technologię z empatią, efektywność z personalizacją oraz automatyzację z autentycznym dialogiem. Firmy, które potrafią właściwie zintegrować AI z kulturą organizacyjną i procesami komunikacji, budują trwałe relacje z klientami oparte na zaufaniu, przewidywalności i wysokiej jakości doświadczeń (Exacto, 2024; EY Polska, 2023; Marszycki, 2020).

Wprowadzenie sztucznej inteligencji prowadzi do głębokiej transformacji w sposobie, w jaki konsumenci podejmują decyzje zakupowe i wchodzi w interakcje z markami. Jak zauważa Lenicka (2025), tradycyjny model liniowej ścieżki klienta, w którym proces zakupowy przebiegał etapami – od świadomości, przez rozważanie, aż po decyzję i lojalność – stopniowo traci na znaczeniu. W jego miejsce pojawia się nieliniowa, dynamiczna sieć interakcji, w której każdy punkt kontaktu z marką może być zarówno początkiem, jak i kontynuacją relacji. Klienci nie poruszają się już po przewidywalnej, zaplanowanej ścieżce, lecz przemieszczają się pomiędzy kanałami, urządzeniami i kontekstami w sposób płynny i często nieprzewidywalny. W tym nowym ekosystemie sztuczna inteligencja pełni funkcję inteligentnego koordynatora doświadczeń klienta, który analizuje dane w czasie rzeczywistym i dostarcza odbiorcy odpowiedni komunikat, we właściwym kanale i w najbardziej dogodnym momencie. Systemy oparte na AI są zdolne do rozpoznawania intencji użytkownika jeszcze zanim ten w pełni uświadomi sobie swoją potrzebę. Dzięki temu marketing ewoluuje od modelu kampanii – opartego na okresowych, masowych działaniach – do modelu ciągłego, zautomatyzowanego dialogu, który utrzymuje relację z klientem w sposób nieprzerwany i kontekstowo dopasowany. Nowy paradygmat relacji z klientem oznacza również, że wartość dla konsumenta nie wynika już wyłącznie z produktu czy usługi,

lecz z całokształtu doświadczenia interakcji z marką. Każdy kontakt – od pierwszej ekspozycji na treść, przez rekomendację produktu, po proces zakupu i obsługi posprzedażowej – współtworzy unikalne doświadczenie, które staje się głównym źródłem przewagi konkurencyjnej. W warunkach rosnącej homogenizacji ofert rynkowych i łatwości kopiowania rozwiązań produktowych, doświadczenie klienta staje się najtrudniejszym do naśladowania zasobem organizacyjnym. Rola AI w tym procesie nie sprowadza się jedynie do automatyzacji kontaktu, ale obejmuje całościową optymalizację podróży klienta. Dzięki integracji danych pochodzących z różnych kanałów komunikacji – stron internetowych, aplikacji mobilnych, mediów społecznościowych czy punktów sprzedaży – sztuczna inteligencja tworzy spójny obraz konsumenta i umożliwia dostosowanie doświadczenia w czasie rzeczywistym. W efekcie przedsiębiorstwa mogą projektować interaktywne, kontekstowe i zindywidualizowane ścieżki zakupowe, które zwiększają satysfakcję, zaangażowanie i długofalową lojalność klientów. Co istotne, w nowym modelu relacji konsument staje się współtwórcą wartości, a nie jedynie odbiorcą przekazu. AI umożliwia identyfikowanie jego preferencji, emocji i reakcji, co pozwala na tworzenie doświadczeń angażujących, spersonalizowanych i emocjonalnie spójnych z tożsamością marki. W rezultacie granica między marketingiem, obsługą a sprzedażą ulega zatarciu, a zarządzanie doświadczeniem klienta przybiera formę ciągłego procesu adaptacji, w którym dane, algorytmy i empatia współdziałają w czasie rzeczywistym. Sztuczna inteligencja redefiniuje współczesny marketing poprzez przesunięcie akcentu z produktu na doświadczenie. Przedsiębiorstwa nie konkurują już wyłącznie na poziomie innowacji technologicznych czy cen, lecz przede wszystkim na poziomie zdolności do projektowania, przewidywania i optymalizacji całościowej podróży klienta. AI staje się kluczową technologią umożliwiającą realizację tej strategii w skali wcześniej niemożliwej – tworząc nowe standardy personalizacji, płynności i spójności doświadczeń, które decydują o sukcesie organizacji w erze cyfrowej (Climenta et al., 2024; Musiał, 2019).

Implementacja sztucznej inteligencji w działaniach marketingowych jest jednym z najpotężniejszych czynników prowadzących do radykalnej poprawy efektywności operacyjnej oraz znaczącego wzrostu zwrotu z inwestycji (ROI). Poprzez inteligentną automatyzację powtarzalnych, czasochłonnych zadań oraz zaawansowane wsparcie w procesach decyzyjnych, AI pozwala przedsiębiorstwom nie tylko osiągnąć trwałą przewagę kosztową, ale również realokować najcenniejszy zasób – kapitał ludzki – do działań o charakterze strategicznym i kreatywnym.

Kluczowym argumentem przemawiającym za inwestycjami w AI jest ich bezpośredni wpływ na rentowność. Narzędzia analityczne oparte na AI umożliwiają optymalizację wydatków marketingowych poprzez dynamiczne, oparte na danych dostosowywanie kampanii w czasie rzeczywistym. Algorytmy potrafią na bieżąco analizować skuteczność poszczególnych kanałów i kreacji, automatycznie przesuwając budżet w kierunku najbardziej efektywnych działań, co bezpośrednio przekłada się na wzrost przychodów i maksymalizację ROI. Jednakże, aby w pełni zrozumieć wartość generowaną przez AI, konieczne jest przyjęcie wielowymiarowego modelu oceny ROI, który wykracza poza proste metryki finansowe. Taki model, zaproponowany w literaturze, wyróżnia trzy kluczowe kategorie zwrotu z inwestycji: (Ramya et al., 2025; Carmichael, 2025):

1. Mierzalny ROI - obejmuje bezpośrednio, kwantyfikowalne korzyści finansowe, takie jak oszczędności kosztów wynikające z automatyzacji, wzrost sprzedaży dzięki lepszej personalizacji czy redukcja kosztów pozyskania klienta;
2. Strategiczny ROI - dotyczy długoterminowych, trudniejszych do zmierzenia korzyści, które wzmacniają pozycję rynkową firmy, takich jak budowanie przewagi konkurencyjnej, zwiększenie udziału w rynku, poprawa wizerunku marki czy zdolność do szybszego wprowadzania innowacji;
3. ROI z potencjału - koncentruje się na wzroście ogólnej dojrzałości technologicznej i organizacyjnej firmy. Obejmuje rozwój nowych kompetencji w zespole, stworzenie kultury organizacyjnej opartej na danych oraz budowanie wewnętrznych zdolności do dalszego, samodzielnego rozwijania i wdrażania rozwiązań AI (Carmichael, 2025).

Automatyzacja procesów marketingowych napędzana przez sztuczną inteligencję obejmuje obecnie niezwykle szerokie spektrum działań, stanowiąc jeden z kluczowych filarów transformacji cyfrowej w obszarze zarządzania relacjami z klientem oraz komunikacji rynkowej. Zastosowanie AI pozwala na przejście od tradycyjnego, ręcznego planowania i realizacji kampanii do w pełni zautomatyzowanych, samouczących się systemów marketingowych, które potrafią analizować dane, podejmować decyzje i dostosowywać działania w sposób autonomiczny. W obszarze marketingu bezpośredniego, systemy oparte na sztucznej inteligencji potrafią kompleksowo zarządzać kampaniami e-mailowymi – od segmentacji odbiorców, przez generowanie treści, po optymalizację momentu wysyłki. Dzięki wykorzystaniu algorytmów uczenia maszynowego i analityki

predykcyjnej możliwe jest personalizowanie zarówno treści, jak i tonu komunikacji w zależności od indywidualnych zachowań użytkowników. Takie systemy potrafią przewidywać, kiedy odbiorca z największym prawdopodobieństwem otworzy wiadomość lub podejmie interakcję z treścią, co prowadzi do znaczącego wzrostu skuteczności kampanii oraz zaangażowania klientów. W rezultacie AI przekształca komunikację e-mailową z narzędzia masowego w precyzyjny instrument relacyjny, dostosowany do rytmu i preferencji każdego odbiorcy. W segmencie B2B, sztuczna inteligencja rewolucjonizuje proces pozyskiwania i kwalifikacji potencjalnych klientów. Algorytmy tzw. lead scoringu analizują dziesiątki, a nawet setki sygnałów – od aktywności w kanałach cyfrowych, przez interakcje z treściami, po dane demograficzne i behawioralne – aby precyzyjnie ocenić prawdopodobieństwo konwersji danego kontaktu. Automatyzacja tego procesu pozwala nie tylko znacząco skrócić cykl sprzedaży, ale również zwiększyć efektywność pracy zespołów handlowych, które mogą skoncentrować swoje działania na najbardziej rokujących leadach. W efekcie dochodzi do lepszego wykorzystania zasobów, redukcji kosztów operacyjnych i podniesienia jakości relacji z klientem biznesowym. Równie dynamicznie rozwija się zastosowanie AI w obszarze reklamy cyfrowej, w szczególności w ramach systemów licytacji w czasie rzeczywistym. Sztuczna inteligencja umożliwia automatyczną analizę milionów parametrów aukcji reklamowych w ułamku sekundy, co pozwala na wybór najbardziej opłacalnych i dopasowanych do odbiorcy ekspozycji reklamowych. Wykorzystanie modeli predykcyjnych i uczenia głębokiego prowadzi do znacznej poprawy efektywności kosztowej kampanii oraz precyzji targetowania – reklamy trafiają do osób o największym potencjale zakupowym, w najbardziej odpowiednim kontekście i czasie. Sztuczna inteligencja nie tylko usprawnia istniejące procesy, ale też umożliwia ich ciągłą optymalizację i samodoskonalenie. Zintegrowane systemy automatyzacji marketingu uczą się na podstawie efektów wcześniejszych działań, co pozwala na iteracyjne ulepszanie komunikacji, strategii targetowania i doboru kanałów. Dzięki temu AI przekształca tradycyjny marketing z procesu reaktywnego w system predykcyjnego uczenia się o rynku, zdolny do bieżącego dostosowywania się do zmieniających się zachowań konsumentów i warunków otoczenia. Automatyzacja procesów marketingowych wspierana przez sztuczną inteligencję prowadzi do jakościowego przełomu w sposobie, w jaki przedsiębiorstwa planują, realizują i analizują swoje działania promocyjne. AI umożliwia masową personalizację, inteligentne zarządzanie cyklem sprzedaży oraz maksymalizację efektywności inwestycji marketingowych. W rezultacie organizacje, które potrafią skutecznie zintegrować te technologie z procesami biznesowymi, zyskują zdolność do

prowadzenia komunikacji precyzyjnej, skalowalnej i adaptacyjnej – stanowiącej kluczowy element przewagi konkurencyjnej w gospodarce cyfrowej (Kapuściński, 2025; Elad, 2025) (tab. 8).

Tabela 8. Kategorie zwrotu z inwestycji (ROI) wynikające z zastosowania sztucznej inteligencji

Kategoria ROI	Kategoria użycia AI	Cele	Oczekiwane korzyści	Metryki, ramy czasowe i łagodzenie ryzyka
Mierzalny zwrot z inwestycji (ROI)	Oparty na sztucznej inteligencji system zarządzania zapasami wykorzystuje algorytmy w celu przewidywania optymalnych poziomów zapasów, automatyzacji procesów zamawiania i minimalizacji przypadków nadwyżek lub niedoborów magazynowych.	Optymalizacja efektywności łańcucha dostaw.	- Obniżenie kosztów utrzymania zapasów - Zmniejszona utrata sprzedaży z powodu braku towarów w magazynie - Poprawa satysfakcji klienta	- 5% kosztów utrzymania i 10% wzrostu sprzedaży w ciągu 1–2 lat - Narzędzie ograniczające ryzyko: audyt i udoskonalanie algorytmów, integracja informacji zwrotnych w czasie rzeczywistym
Strategiczny zwrot z inwestycji	System zarządzania zapasami oparty na sztucznej inteligencji umożliwia śledzenie zapasów w czasie rzeczywistym i dostosowywanie zapasów na podstawie trendów popytu rynkowego i wahań sezonowych.	Poprawa obsługi klienta i lojalności dzięki niezawodnej dostępności produktów.	- Usprawnione procesy inwentaryzacyjne - Szybka reakcja na trendy rynkowe	- 20% redukcji nieefektywności procesów w ciągu 3–5 lat - O 15% więcej zakupów od powracających klientów w ciągu najbliższych 3 lat - Narzędzie ograniczające ryzyko: wdrożenie zarządzania danymi i ciągłego szkolenia modeli AI w oparciu o różnicowane i aktualne źródła danych w celu zapewnienia dokładności i niezawodności
Zwrot z inwestycji w możliwości	Zwiększanie możliwości technologicznych poprzez podnoszenie kompetencji pracowników w zakresie sztucznej inteligencji, zapewnianie szkoleń z zakresu systemu zarządzania zapasami i rekrutacja na stanowiska wspierające integrację sztucznej inteligencji w organizacji.	Rozwijanie umiejętności pracowników i kompetencji technologicznych, aby stać się organizacją opartą na sztucznej inteligencji.	- Zwiększona biegłość pracowników w zakresie systemów AI w celu zwiększenia wydajności - Wzmocniona kultura innowacji	- Ciągła poprawa wydajności pracowników - Czynniki ograniczające ryzyko: zapewnienie programów kształcenia ustawicznego pozwalających na dostosowanie się do zmian technologicznych

Źródło: Opracowanie własne na podstawie: Carmichael, M. (2025). Demonstrating AI ROI: How to Measure and Prove the Value of Your AI Investments. ISACA. <https://www.isaca.org/resources/news-and-trends/newsletters/atisaca/2025/volume-5/how-to-measure-and-prove-the-value-of-your-ai-investments>

W tym kontekście kluczowa staje się koncepcja „wzmocnionego pracownika”. Wbrew obawom o masowe zastępowanie ludzi przez maszyny, dominującym modelem współpracy staje się synergia między inteligencją ludzką a sztuczną. AI nie tylko przejmuje powtarzalne, mechaniczne zadania (mechaniczna AI), ale przede wszystkim wzmacnia zdolności analityczne i decyzyjne człowieka (myśląca AI). Marketerzy, wyposażeni w narzędzia AI, mogą w ciągu kilku minut analizować zbiory danych, których ręczne przetworzenie zajęłoby tygodnie. Otrzymują oni gotowe rekomendacje, identyfikacje trendów i prognozy, co pozwala im skupić się na strategicznym planowaniu, kreatywnym myśleniu i budowaniu relacji z klientami. AI staje się inteligentnym asystentem i partnerem w procesie analitycznym, a nie tylko wykonawcą poleceń (Plangger & Campbell, 2022).

Badania empiryczne nad wdrażaniem AI w przedsiębiorstwach

Metodologia badań ankietowych i charakterystyka próby

W niniejszej pracy przedstawiono szczegółowy opis przebiegu badań ankietowych oraz charakterystykę próby badawczej, która stanowiła podstawę analiz statystycznych. Punktem wyjścia była analiza danych pochodzących ze 177 przedsiębiorstw, obejmujących szeroki zakres zmiennych opisujących zarówno cechy organizacyjne, jak i poziom wdrożenia oraz integracji sztucznej inteligencji. Uwzględniono ponadto czynniki dotyczące proaktywności, działań marketingowych, korzystania z narzędzi AI oraz kluczowych wyników biznesowych, takich jak satysfakcja klientów, rentowność i innowacyjność.

Z uwagi na złożoność zjawiska wdrażania AI w przedsiębiorstwach, zasadniczym celem przeprowadzonych analiz było określenie, które czynniki organizacyjne i technologiczne w sposób istotny wpływają na efektywność działalności firm. W konsekwencji skupiono się także na identyfikacji zależności statystycznych pomiędzy wybranymi zmiennymi, pozwalających określić kierunek oraz siłę tych relacji.

W niniejszej pracy sformułowano pięć hipotez badawczych:

- **(H1)** Poziom wdrożenia sztucznej inteligencji (indeks AI) koreluje z wynikami biznesowymi, w tym z satysfakcją klientów, rentownością oraz innowacyjnością;
- **(H2)** Przedsiębiorstwa z branży technologicznej charakteryzują się wyższym poziomem zaawansowania AI niż przedsiębiorstwa działające w innych sektorach gospodarki;

- **(H3)** Stosowanie wybranych narzędzi AI (np. Mailchimp) oddziałuje na wyniki marketingowe oraz rentowność firm;
- **(H4)** Wielkość przedsiębiorstwa moderuje relację pomiędzy wdrożeniem AI a osiąganymi wynikami biznesowymi;
- **(H5)** Występuje związek między zaawansowanymi praktykami marketingowymi (takimi jak alokacja zasobów czy monitorowanie potrzeb klientów) a wynikami biznesowymi przedsiębiorstw.

Sformułowanie powyższych hipotez pozwoliło zaplanować badania w sposób umożliwiający empiryczne potwierdzenie lub odrzucenie przyjętych założeń. W konsekwencji przeprowadzone badania empiryczne miały na celu rozpoznanie poziomu wdrożenia oraz zakresu wykorzystania technologii sztucznej inteligencji w przedsiębiorstwach działających w Polsce, ze szczególnym uwzględnieniem podmiotów funkcjonujących w ramach Specjalnych Stref Ekonomicznych. Zastosowano podejście ilościowe oparte na metodzie sondażu diagnostycznego, przy wykorzystaniu kwestionariusza ankiety, co umożliwiło zebranie porównywalnych danych ilościowych i ich dalszą analizę statystyczną.

Na podstawie zdefiniowanych celów badawczych opracowano autorski kwestionariusz ankiety, który powstał w wyniku analizy literatury dotyczącej zastosowań sztucznej inteligencji w zarządzaniu, marketingu i strategiach organizacyjnych. Kwestionariusz przygotowano w formie elektronicznej (formularz online), co umożliwiło szybkie dotarcie do respondentów z różnych branż i regionów kraju. Takie rozwiązanie przyczyniło się do zwiększenia dostępności badania, a jednocześnie pozwoliło ograniczyć jego koszty oraz skrócić czas realizacji projektu.

Dobór próby miał charakter celowy i obejmował osoby zajmujące stanowiska kierownicze, specjalistyczne lub właścicielskie w przedsiębiorstwach reprezentujących różne sektory gospodarki. Respondenci zostali poproszeni o ocenę stopnia wdrożenia i wykorzystania AI w swoich organizacjach, z uwzględnieniem takich obszarów, jak infrastruktura technologiczna, integracja z planowaniem strategicznym, działania marketingowe oraz ogólne wyniki organizacyjne.

Na podstawie przyjętej procedury badawczej ankietę przeprowadzono w okresie od marca do maja 2025 roku. Zgromadzony materiał opracowano z zastosowaniem metod statystyki opisowej, obejmujących obliczenie częstości i udziałów procentowych, co umożliwiło uzyskanie pełnego obrazu struktury badanej próby.

W badaniu uczestniczyło 177 respondentów reprezentujących przedsiębiorstwa zróżnicowane pod względem wielkości, branży oraz poziomu zaawansowania

we wdrażaniu technologii AI. W strukturze wielkości dominowały firmy małe (29,40%) i średnie (29,40%), podczas gdy mikroprzedsiębiorstwa stanowiły 22,00%, a duże organizacje – 19,20%. Z kolei analiza struktury branżowej wykazała przewagę przedsiębiorstw z sektora usług (35,00%) i handlu (30,50%), przy mniejszym udziale podmiotów z branż produkcyjnej (15,80%), technologicznej (11,90%) oraz pozostałych (6,80%). Takie zróżnicowanie próby zapewniło reprezentatywność danych oraz pozwoliło na identyfikację różnic w podejściu do wdrażania AI pomiędzy poszczególnymi typami przedsiębiorstw.

Wyniki analizy poziomu zaawansowania wdrożenia sztucznej inteligencji w badanych przedsiębiorstwach wskazują, że większość respondentów oceniła go jako raczej niski (33,90%) lub średni (28,20%). Stosunkowo niewielki odsetek uczestników badania – jedynie 11,90% – określił poziom wdrożenia jako raczej zaawansowany lub bardzo zaawansowany. Uzyskane wyniki potwierdzają, że większość firm wciąż znajduje się na wczesnym etapie adaptacji rozwiązań opartych na sztucznej inteligencji, co może wynikać z ograniczonych zasobów technologicznych lub braku specjalistycznych kompetencji.

Uwzględniając terytorialne zróżnicowanie badanych podmiotów, zauważono, że dominowały organizacje z miast średnich (56,50%) oraz małych (30,50%). Udział przedsiębiorstw z dużych miast wyniósł 8,50%, a z obszarów wiejskich – 4,50%. Struktura ta sugeruje, że wdrażanie technologii AI nie jest zjawiskiem ograniczonym wyłącznie do dużych ośrodków miejskich, lecz coraz częściej obejmuje również mniejsze miejscowości, w których przedsiębiorstwa poszukują narzędzi usprawniających procesy biznesowe.

Analizując doświadczenie zawodowe respondentów, można zauważyć, że największą grupę stanowili pracownicy zatrudnieni od 1 do 3 lat (43,50%) oraz krócej niż 1 rok (32,20%). Mniejsze udziały miały osoby pracujące od 4 do 6 lat (13,00%) oraz powyżej 6 lat (11,30%). Wyniki te mogą świadczyć o relatywnie wysokiej rotacji pracowników w badanych przedsiębiorstwach, co potencjalnie utrudnia stabilne wdrażanie i rozwój projektów z zakresu sztucznej inteligencji.

W kontekście struktury stanowisk najliczniejszą grupę respondentów stanowili specjaliści (36,20%), co wskazuje na znaczący udział personelu operacyjnego w procesach implementacji AI. Kolejnymi grupami byli dyrektorzy (19,80%) oraz menedżerowie (19,20%), natomiast mniejszy udział odnotowano wśród właścicieli lub prezesów (11,90%) oraz osób pełniących inne funkcje (13,00%). Takie proporcje mogą sugerować, że inicjatywy związane z wykorzystaniem AI są w dużej mierze realizowane na poziomie specjalistycznym, przy ograniczonym bezpośrednim zaangażowaniu najwyższej kadry zarządzającej (tab. 9).

Tabela 9. Charakterystyka badanej próby przedsiębiorstw

Kategoria	Podkategoria	N	%
Wielkość organizacji	Duże	34	19,20%
	Małe	52	29,40%
	Mikro	39	22,00%
	Średnie	52	29,40%
	Łącznie	177	100,00%
Branża	Handel	54	30,50%
	Inne	12	6,80%
	Produkcja	28	15,80%
	Technologie	21	11,90%
	Usługi	62	35,00%
	Łącznie	177	100,00%
Ogólny poziom AI w organizacji	Bardzo niski	46	26,00%
	Raczej niski	60	33,90%
	Średni	50	28,20%
	Raczej zaawansowany	12	6,80%
	Bardzo zaawansowany	9	5,10%
	Łącznie	177	100,00%
Lokalizacja organizacji	Duże miasto	15	8,50%
	Małe miasto	54	30,50%
	Wieś	8	4,50%
	Średnie miasto	100	56,50%
	Łącznie	177	100,00%
Staż pracy w obecnej organizacji	<1 rok	57	32,20%
	1–3 lata	77	43,50%
	4–6 lat	23	13,00%
	>6 lat	20	11,30%
	Łącznie	177	100,00%
Stanowisko respondenta	Dyrektor	35	19,80%
	Inne	23	13,00%
	Manager	34	19,20%
	Specjalista	64	36,20%
	Właściciel/CEO	21	11,90%
	Łącznie	177	100,00%

Źródło: Opracowanie własne.

Podsumowując, próba badawcza obejmowała zróżnicowaną grupę przedsiębiorstw, wśród których dominowały małe i średnie podmioty z sektorów usługowego oraz handlowego. Zgromadzone dane wskazują, że większość z nich znajduje się dopiero na początkowym etapie wdrażania rozwiązań opartych na sztucznej inteligencji, co potwierdza potrzebę dalszego wsparcia transformacji cyfrowej w tym segmencie przedsiębiorstw.

Zastosowana metodologia oraz struktura próby umożliwiły przeprowadzenie kompleksowej analizy poziomu zaawansowania wdrożeń sztucznej inteligencji w różnych obszarach działalności przedsiębiorstw, zapewniając jednocześnie wysoki poziom wiarygodności i reprezentatywności uzyskanych wyników.

Analiza wyników badań (infrastruktura, strategia, marketing, wyniki organizacyjne)

W niniejszym podrozdziale zaprezentowano szczegółowe wyniki badań empirycznych dotyczących poziomu zaawansowania wdrożeń sztucznej inteligencji w przedsiębiorstwach w czterech kluczowych obszarach: infrastruktury technologicznej, integracji strategicznej, zastosowań marketingowych oraz osiąganych wyników organizacyjnych. Analiza obejmuje zarówno aspekty techniczne, takie jak dostępność usług chmurowych, aplikacji i architektury danych dla AI, jak i elementy zarządzania strategicznego, w tym stopień włączenia rozwiązań sztucznej inteligencji w proces planowania biznesowego. W dalszej części przedstawiono wyniki dotyczące praktyk marketingowych wykorzystujących narzędzia AI do analizy rynku, prognozowania trendów oraz personalizacji działań promocyjnych. Ostatnia część analizy koncentruje się na ocenie wpływu wdrożenia AI na rezultaty organizacyjne, takie jak rentowność, satysfakcja klientów czy innowacyjność przedsiębiorstw.

Na podstawie danych zaprezentowanych w tabeli 10 można zauważyć, że 49,20% respondentów „zdecydowanie nie” zgadza się ze stwierdzeniem, iż ich organizacja posiada usługi zarządzania danymi i architektury dla AI, a 11,90% odpowiedziało „raczej nie”. Odpowiedź neutralną („ani tak, ani nie”) wskazało 20,30% osób, natomiast 7,90% respondentów zaznaczyło „raczej tak”, a 10,70% – „zdecydowanie tak”. W odniesieniu do rozwiniętych usług komunikacji sieciowej i chmurowej dla AI, 43,50% badanych odpowiedziało „zdecydowanie nie”, 16,90% – „raczej nie”, 14,70% – „ani tak, ani nie”, 12,40% – „raczej tak”, a kolejne 12,40% – „zdecydowanie tak”. Na stwierdzenie dotyczące posiadania portfela aplikacji i usług AI, takich jak Microsoft Cognitive Services

czy Google Cloud Vision, 40,70% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 18,10% – „raczej nie”, 20,90% – „ani tak, ani nie”, 10,70% – „raczej tak”, a 9,60% – „zdecydowanie tak”. W przypadku infrastruktury obejmującej operacje i funkcje AI, takie jak serwery, procesory o dużej mocy czy monitory wydajności, 44,60% uczestników badania wskazało „zdecydowanie nie”, 17,50% – „raczej nie”, 15,80% – „ani tak, ani nie”, 7,90% – „raczej tak” i 14,10% – „zdecydowanie tak”. W odniesieniu do zapewniania bezpieczeństwa danych przez infrastrukturę AI, 36,20% respondentów wybrało „zdecydowanie nie”, 18,10% – „raczej nie”, 17,50% – „ani tak, ani nie”, 11,90% – „raczej tak”, a 16,40% – „zdecydowanie tak” (tab. 10).

Tabela 10. Kompetencje infrastruktury AI

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Nasza organizacja posiada usługi zarządzania danymi i architektury dla AI	49,20%	11,90%	20,30%	7,90%	10,70%
Mamy rozwinięte usługi komunikacji sieciowej i chmurowe dla AI	43,50%	16,90%	14,70%	12,40%	12,40%
Dysponujemy portfelem aplikacji i usług AI (np. Microsoft Cognitive Services, Google Cloud Vision)	40,70%	18,10%	20,90%	10,70%	9,60%
Nasza infrastruktura AI obejmuje operacje/ funkcje, takie jak serwery, procesory o dużej mocy czy monitory wydajności	44,60%	17,50%	15,80%	7,90%	14,10%
Nasza infrastruktura AI zapewnia bezpieczeństwo danych	36,20%	18,10%	17,50%	11,90%	16,40%

Źródło: Opracowanie własne.

Analizując dane przedstawione w tabeli 11, można zauważyć, w jakim stopniu badane organizacje integrują zagadnienia związane ze sztuczną inteligencją z procesem planowania strategicznego i biznesowego. W przypadku stwierdzenia „Posiadamy wizję, jak AI przyczynia się do zwiększenia wartości biznesowej

naszej organizacji”, 32,80% respondentów odpowiedziało „zdecydowanie nie”, 13,00% – „raczej nie”, 18,60% – „ani tak, ani nie”, 12,40% – „raczej tak”, a 23,20% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „Plany strategiczne organizacji są zintegrowane z planowaniem w zakresie AI”, 35,60% badanych wskazało odpowiedź „zdecydowanie nie”, 16,40% – „raczej nie”, 15,30% – „ani tak, ani nie”, 9,60% – „raczej tak”, a 23,20% – „zdecydowanie tak”. Na pytanie „Zarząd na poziomie funkcjonalnym i strategicznym wykazuje zrozumienie w inwestycji w AI”, 33,90% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 13,00% – „raczej nie”, 15,80% – „ani tak, ani nie”, 13,00% – „raczej tak”, a 24,30% – „zdecydowanie tak”. W przypadku stwierdzenia „Nasze plany biznesowe uwzględniają ciągłą inwestycję w AI”, 39,00% badanych wybrało „zdecydowanie nie”, 12,40% – „raczej nie”, 18,10% – „ani tak, ani nie”, 12,40% – „raczej tak”, a 18,10% – „zdecydowanie tak” (tab. 11).

Tabela 11. Integracja AI z planowaniem biznesowymi

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Posiadamy wizję, jak AI przyczynia się do zwiększenia wartości biznesowej naszej organizacji	32,80%	13,00%	18,60%	12,40%	23,20%
Plany strategiczne organizacji są zintegrowane z planowaniem w zakresie AI	35,60%	16,40%	15,30%	9,60%	23,20%
Zarząd na poziomie funkcjonalnym i strategicznym wykazuje zrozumienie w inwestycji w AI	33,90%	13,00%	15,80%	13,00%	24,30%
Nasze plany biznesowe uwzględniają ciągłą inwestycję w AI	39,00%	12,40%	18,10%	12,40%	18,10%

Źródło: Opracowanie własne.

W tabeli 12 przedstawiono wyniki obrazujące, w jakim zakresie badane przedsiębiorstwa podejmują inicjatywy ukierunkowane na proaktywne wykorzystanie sztucznej inteligencji. Z tabeli wynika, że w odniesieniu do stwierdzenia „Regularnie eksperymentujemy z nowymi narzędziami i technikami AI”, 29,90% respondentów wskazało odpowiedź „zdecydowanie nie”, 14,70% – „raczej nie”,

18,60% – „ani tak, ani nie”, 16,40% – „raczej tak”, a 20,30% – „zdecydowanie tak”. W przypadku stwierdzenia „Kultura naszej organizacji wspiera stosowanie nowych sposobów wykorzystania AI”, 35,00% badanych wybrało „zdecydowanie nie”, 16,90% – „raczej nie”, 15,30% – „ani tak, ani nie”, 12,40% – „raczej tak”, a 20,30% – „zdecydowanie tak”. Dla stwierdzenia „Stale poszukujemy nowych sposobów zwiększania efektywności wykorzystania AI”, 31,60% respondentów odpowiedziało „zdecydowanie nie”, 16,40% – „raczej nie”, 15,80% – „ani tak, ani nie”, 15,80% – „raczej tak”, natomiast 20,30% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „Poszukujemy innowacji w zakresie AI w celu zwiększenia satysfakcji i lojalności klientów”, 31,60% uczestników badania wskazało „zdecydowanie nie”, 18,60% – „raczej nie”, 11,90% – „ani tak, ani nie”, 13,00% – „raczej tak”, a 24,90% – „zdecydowanie tak”.

Tabela 12. Proaktywne podejście do AI

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Regularnie eksperymentujemy z nowymi narzędziami i technikami AI	29,90%	14,70%	18,60%	16,40%	20,30%
Kultura naszej organizacji wspiera stosowanie nowych sposobów wykorzystania AI	35,00%	16,90%	15,30%	12,40%	20,30%
Stale poszukujemy nowych sposobów zwiększania efektywności wykorzystania AI	31,60%	16,40%	15,80%	15,80%	20,30%
Poszukujemy innowacji w zakresie AI w celu zwiększenia satysfakcji i lojalności klientów	31,60%	18,60%	11,90%	13,00%	24,90%

Źródło: Opracowanie własne.

Kolejną kwestią poddaną analizie, przedstawioną w tabeli 13, jest zakres wykorzystania narzędzi sztucznej inteligencji w zarządzaniu informacją marketingową przez badane przedsiębiorstwa. Z danych przedstawionych w tabeli 13 wynika, że dla stwierdzenia „Wykorzystujemy narzędzia AI do gromadzenia informacji o klientach” 27,70% respondentów wybrało odpowiedź „zdecydowanie nie”, 14,10% – „raczej nie”, 19,20% – „ani tak, ani nie”, 16,90% – „raczej tak”, a 22,00% – „zdecydowanie tak”. W przypadku stwierdzenia „Wykorzystujemy narzędzia AI do gromadzenia informacji o konkurentach”, 33,90%

badanych wskazało „zdecydowanie nie”, 14,70% – „raczej nie”, 13,00% – „ani tak, ani nie”, 14,70% – „raczej tak”, a 23,70% – „zdecydowanie tak”. Dla stwierdzenia „Wspieramy rozwój programów marketingowych za pomocą AI do analizy danych rynkowych”, 31,10% respondentów odpowiedziało „zdecydowanie nie”, 16,40% – „raczej nie”, 15,80% – „ani tak, ani nie”, 15,30% – „raczej tak”, natomiast 21,50% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „Śledzimy potrzeby i oczekiwania klientów w czasie rzeczywistym za pomocą algorytmów AI dla zwiększenia ich satysfakcji”, 39,00% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 12,40% – „raczej nie”, 13,00% – „ani tak, ani nie”, 11,90% – „raczej tak”, a 23,70% – „zdecydowanie tak”. Na stwierdzenie „Wykorzystujemy generowane przez AI raporty i analizy marketingowe”, 29,40% badanych wskazało „zdecydowanie nie”, 14,10% – „raczej nie”, 18,60% – „ani tak, ani nie”, 14,70% – „raczej tak”, a 23,20% – „zdecydowanie tak”. W przypadku stwierdzenia „Stosujemy AI do prognozowania trendów rynkowych i zachowań klientów takich jak zaufanie i zaangażowanie”, 37,30% respondentów wybrało „zdecydowanie nie”, 16,90% – „raczej nie”, 9,60% – „ani tak, ani nie”, 16,40% – „raczej tak”, a 19,80% – „zdecydowanie tak”.

Tabela 13. Marketing Information Management z zastosowaniem AI

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Wykorzystujemy narzędzia AI do gromadzenia informacji o klientach	27,70%	14,10%	19,20%	16,90%	22,00%
Wykorzystujemy narzędzia AI do gromadzenia informacji o konkurentach	33,90%	14,70%	13,00%	14,70%	23,70%
Wspieramy rozwój programów marketingowych za pomocą AI do analizy danych rynkowych	31,10%	16,40%	15,80%	15,30%	21,50%
Śledzimy potrzeby i oczekiwania klientów w czasie rzeczywistym za pomocą algorytmów AI dla zwiększenia ich satysfakcji	39,00%	12,40%	13,00%	11,90%	23,70%

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Wykorzystujemy generowane przez AI raporty i analizy marketingowe	29,40%	14,10%	18,60%	14,70%	23,20%
Stosujemy AI do prognozowania trendów rynkowych i zachowań klientów takich jak zaufanie i zaangażowanie	37,30%	16,90%	9,60%	16,40%	19,80%

Źródło: Opracowanie własne.

Następnym obszarem poddanym analizie, zaprezentowanym w tabeli 14, jest wykorzystanie sztucznej inteligencji w procesie planowania działań marketingowych. Z tabeli wynika, że w przypadku stwierdzenia „Wykorzystujemy algorytmy AI do segmentowania i targetowania rynku”, 33,30% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 14,70% – „raczej nie”, 19,20% – „ani tak, ani nie”, 14,70% – „raczej tak”, a 18,10% – „zdecydowanie tak”. Dla stwierdzenia „Wdrażamy narzędzia AI do opracowywania personalizowanych strategii marketingowych”, 39,50% badanych wskazało odpowiedź „zdecydowanie nie”, 19,20% – „raczej nie”, 9,00% – „ani tak, ani nie”, 13,60% – „raczej tak”, a 18,60% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „Nasze procesy planowania marketingowego są wspierane przez modele predykcyjne AI”, 38,40% respondentów wybrało „zdecydowanie nie”, 18,60% – „raczej nie”, 17,50% – „ani tak, ani nie”, 7,90% – „raczej tak”, natomiast 17,50% – „zdecydowanie tak”. W przypadku stwierdzenia „Stosujemy AI do oceny skuteczności różnych scenariuszy marketingowych jeszcze przed ich wdrożeniem”, 27,10% uczestników badania wskazało „zdecydowanie nie”, 13,60% – „raczej nie”, 17,50% – „ani tak, ani nie”, 15,30% – „raczej tak”, a 26,60% – „zdecydowanie tak”. Dla stwierdzenia „AI wspomaga nasze decyzje dotyczące alokacji budżetów marketingowych na różne kanały komunikacji”, 29,40% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 14,70% – „raczej nie”, 16,40% – „ani tak, ani nie”, 14,70% – „raczej tak”, oraz 24,90% – „zdecydowanie tak”.

Tabela 14. Planowanie marketingowe z zastosowaniem AI

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Wykorzystujemy algorytmy AI do segmentowania i targetowania rynku	33,30%	14,70%	19,20%	14,70%	18,10%
Wdrażamy narzędzia AI do opracowywania personalizowanych strategii marketingowych	39,50%	19,20%	9,00%	13,60%	18,60%
Nasze procesy planowania marketingowego są wspierane przez modele predykcyjne AI	38,40%	18,60%	17,50%	7,90%	17,50%
Stosujemy AI do oceny skuteczności różnych scenariuszy marketingowych jeszcze przed ich wdrożeniem	27,10%	13,60%	17,50%	15,30%	26,60%
AI wspomaga nasze decyzje dotyczące alokacji budżetów marketingowych na różne kanały komunikacji	29,40%	14,70%	16,40%	14,70%	24,90%

Źródło: Opracowanie własne.

Kolejnym zagadnieniem przedstawionym w tabeli 15 jest zastosowanie sztucznej inteligencji w procesie realizacji i organizacji działań marketingowych. Z danych przedstawionych w tabeli wynika, że w przypadku stwierdzenia „Optymalizujemy alokację zasobów marketingowych za pomocą AI”, 29,90% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 10,70% – „raczej nie”, 19,80% – „ani tak, ani nie”, 15,30% – „raczej tak”, a 24,30% – „zdecydowanie tak”. Dla stwierdzenia „Korzystamy z AI do automatyzacji organizacji i realizacji kampanii marketingowych”, 31,60% badanych wskazało odpowiedź „zdecydowanie nie”, 15,80% – „raczej nie”, 13,00% – „ani tak, ani nie”, 21,50% – „raczej tak”, a 18,10% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „Wykorzystujemy systemy AI do dynamicznego dostosowywania działań marketingowych w odpowiedzi na zmieniające się warunki rynkowe”, 33,90% respondentów

wybrało „zdecydowanie nie”, 18,10% – „raczej nie”, 13,60% – „ani tak, ani nie”, 13,60% – „raczej tak”, a 20,90% – „zdecydowanie tak”. W przypadku stwierdzenia „Zastosowanie AI wspiera monitorowanie i analizę wyników kampanii marketingowych w czasie rzeczywistym”, 30,50% uczestników badania wskazało „zdecydowanie nie”, 13,00% – „raczej nie”, 15,80% – „ani tak, ani nie”, 15,80% – „raczej tak”, a 24,90% – „zdecydowanie tak”. Dla stwierdzenia „Generujemy treści marketingowe przy użyciu narzędzi AI, co zwiększa ich jakość i efektywność”, 37,90% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 9,60% – „raczej nie”, 15,80% – „ani tak, ani nie”, 10,20% – „raczej tak”, oraz 26,60% – „zdecydowanie tak”.

Tabela 15. Implementacja marketingowa z zastosowaniem AI

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
Optymalizujemy alokację zasobów marketingowych za pomocą AI	29,90%	10,70%	19,80%	15,30%	24,30%
Korzystamy z AI do automatyzacji organizacji i realizacji kampanii marketingowych	31,60%	15,80%	13,00%	21,50%	18,10%
Wykorzystujemy systemy AI do dynamicznego dostosowywania działań marketingowych w odpowiedzi na zmieniające się warunki rynkowe	33,90%	18,10%	13,60%	13,60%	20,90%
Zastosowanie AI wspiera monitorowanie i analizę wyników kampanii marketingowych w czasie rzeczywistym	30,50%	13,00%	15,80%	15,80%	24,90%
Generujemy treści marketingowe przy użyciu narzędzi AI, co zwiększa ich jakość i efektywność	37,90%	9,60%	15,80%	10,20%	26,60%

Źródło: Opracowanie własne.

Obszarem zaprezentowanym w tabeli 16 są wyniki organizacyjne badanych przedsiębiorstw w porównaniu z ich głównymi konkurentami. Z danych przedstawionych w tabeli wynika, że w przypadku stwierdzenia „W porównaniu do głównych konkurentów nasza organizacja kreuje wartości dla klienta”, 26,00% respondentów udzieliło odpowiedzi „zdecydowanie nie”, 14,10% – „raczej nie”, 18,10% – „ani tak, ani nie”, 18,60% – „raczej tak”, a 23,20% – „zdecydowanie tak”. Dla stwierdzenia „W porównaniu do głównych konkurentów mamy więcej stałych klientów”, 32,20% badanych wybrało „zdecydowanie nie”, 13,60% – „raczej nie”, 18,10% – „ani tak, ani nie”, 13,60% – „raczej tak”, a 22,60% – „zdecydowanie tak”. W odniesieniu do stwierdzenia „W porównaniu do głównych konkurentów mamy więcej usatysfakcjonowanych klientów”, 29,40% respondentów wskazało odpowiedź „zdecydowanie nie”, 15,30% – „raczej nie”, 15,30% – „ani tak, ani nie”, 14,70% – „raczej tak”, oraz 25,40% – „zdecydowanie tak”. W przypadku stwierdzenia „W porównaniu do głównych konkurentów jesteśmy bardziej rentowni”, 33,30% uczestników badania udzieliło odpowiedzi „zdecydowanie nie”, 18,10% – „raczej nie”, 13,60% – „ani tak, ani nie”, 18,10% – „raczej tak”, a 16,90% – „zdecydowanie tak”. Dla stwierdzenia „W porównaniu do głównych konkurentów jesteśmy bardziej innowacyjni”, 28,20% respondentów wybrało odpowiedź „zdecydowanie nie”, 17,50% – „raczej nie”, 17,50% – „ani tak, ani nie”, 9,00% – „raczej tak”, a 27,70% – „zdecydowanie tak”.

Tabela 16. Wyniki organizacyjne

	Zdecydowanie nie	Raczej nie	Ani tak, ani nie	Raczej tak	Zdecydowanie tak
W porównaniu do głównych konkurentów nasza organizacja kreuje wartości dla klienta	26,00%	14,10%	18,10%	18,60%	23,20%
W porównaniu do głównych konkurentów mamy więcej stałych klientów	32,20%	13,60%	18,10%	13,60%	22,60%
W porównaniu do głównych konkurentów mamy więcej usatysfakcjonowanych klientów	29,40%	15,30%	15,30%	14,70%	25,40%

W porównaniu do głównych konkurentów jesteśmy bardziej rentowni	33,30%	18,10%	13,60%	18,10%	16,90%
W porównaniu do głównych konkurentów jesteśmy bardziej innowacyjni	28,20%	17,50%	17,50%	9,00%	27,70%

Źródło: Opracowanie własne.

W tabeli 17 zaprezentowano zestawienie narzędzi sztucznej inteligencji wykorzystywanych w komunikacji marketingowej, wraz z odsetkiem organizacji deklarujących ich stosowanie. W przypadku narzędzia HubSpot, 13,00% respondentów wskazało, że jest ono stosowane w ich organizacji, natomiast 87,00% zadeklarowało, że nie jest stosowane. Dla narzędzia Mailchimp odsetek stosowania wyniósł 32,80%, a 67,20% badanych wskazało, że nie wykorzystuje tego rozwiązania. W odniesieniu do Hootsuite, 33,30% respondentów potwierdziło jego stosowanie, a 66,70% – brak wykorzystania. W przypadku narzędzia AdRoll, 16,40% uczestników badania zadeklarowało jego użycie, natomiast 83,60% – że nie jest ono stosowane. Dla Salesforce Einstein 23,70% badanych potwierdziło stosowanie, a 76,30% – brak wykorzystania. W odniesieniu do Pardot, 16,40% respondentów wskazało, że narzędzie jest stosowane, natomiast 83,60% – że nie. Dla Cortex, 33,30% organizacji korzysta z tego rozwiązania, a 66,70% – nie korzysta. W przypadku Conversica, 37,30% badanych potwierdziło stosowanie, natomiast 62,70% – brak wykorzystania. Dla Drift, 18,10% respondentów wskazało, że narzędzie jest stosowane, a 81,90% – że nie. W odniesieniu do Phrasee, 19,80% badanych potwierdziło jego wykorzystanie, natomiast 80,20% – brak stosowania. W przypadku MarketMuse, 27,10% organizacji stosuje to narzędzie, a 72,90% – nie wykorzystuje go w swojej działalności. Dla Crimson Hexagon, 22,60% respondentów potwierdziło stosowanie, natomiast 77,40% – że narzędzie nie jest używane. W odniesieniu do Optmyzr, 18,10% badanych organizacji korzysta z tego rozwiązania, natomiast 81,90% – nie korzysta. W przypadku Emarsys, 29,40% respondentów wskazało, że narzędzie jest stosowane w ich organizacji, a 70,60% – że nie jest wykorzystywane. Dla Undetectable.ai, 13,60% uczestników badania potwierdziło jego użycie, natomiast 86,40% – brak stosowania.

Tabela 17. Narzędzia AI do komunikacji marketingowej stosowane w badanym przedsiębiorstwie

Narzędzie AI	Stosuje się	Nie stosuje się
HubSpot	13,00%	87,00%
Mailchimp	32,80%	67,20%
Hootsuite	33,30%	66,70%
AdRoll	16,40%	83,60%
Salesforce Einstein	23,70%	76,30%
Pardot	16,40%	83,60%
Cortex	33,30%	66,70%
Conversica	37,30%	62,70%
Drift	18,10%	81,90%
Phrasee	19,80%	80,20%
MarketMuse	27,10%	72,90%
Crimson Hexagon	22,60%	77,40%
Optmyzr	18,10%	81,90%
Emarsys	29,40%	70,60%
Undetectable.ai	13,60%	86,40%

Źródło: Opracowanie własne.

Podsumowując, uzyskane wyniki wskazują, że mimo rosnącego zainteresowania sztuczną inteligencją, poziom jej wdrożenia w badanych przedsiębiorstwach pozostaje umiarkowany. Najsłabiej rozwiniętym obszarem jest infrastruktura technologiczna, obejmująca architekturę danych, usługi chmurowe i systemy wspierające bezpieczeństwo informacji. Z kolei nieco wyższy stopień zaawansowania można zaobserwować w sferze działań marketingowych, szczególnie w zakresie wykorzystania AI do analizy danych rynkowych, segmentacji klientów czy automatyzacji kampanii. Wyniki te sugerują, że kluczowym wyzwaniem dla organizacji jest nie tylko inwestowanie w technologię, lecz także jej strategiczna integracja z celami biznesowymi i kulturą innowacji. Ostatecznie skuteczne wykorzystanie AI wymaga spójnego podejścia, łączącego rozwój infrastruktury, kompetencji i procesów organizacyjnych z orientacją na długofalowe wyniki biznesowe.

Weryfikacja hipotez badawczych

W niniejszym podrozdziale przedstawiono proces empirycznej weryfikacji hipotez badawczych sformułowanych na podstawie założeń teoretycznych dotyczących wdrażania sztucznej inteligencji w przedsiębiorstwach. Celem analiz było zidentyfikowanie związków między stopniem zaawansowania technologicznego organizacji a osiąganymi wynikami biznesowymi, obejmującymi takie obszary jak satysfakcja klientów, rentowność oraz innowacyjność. Weryfikacja poszczególnych hipotez opierała się na zastosowaniu metod statystycznych odpowiednich do charakteru zmiennych – w szczególności analizy korelacji i regresji – co pozwoliło ocenić kierunek oraz siłę relacji pomiędzy badanymi zjawiskami.

Hipoteza 1 odnosiła się do zależności występującej pomiędzy poziomem zaawansowania wdrożenia sztucznej inteligencji a wynikami biznesowymi organizacji.

Tabela 18. Wartości współczynników korelacji Pearsona pomiędzy indeksem zaawansowania AI a wynikami biznesowymi

Zmienna niezależna	Zmienna zależna	N	Korelacja (r)	p-wartość
Zaawansowanie planowania marketingowego	Satysfakcja klientów	177	-0,006	0,935
Zaawansowanie planowania marketingowego	Rentowność	177	-0,025	0,736
Zaawansowanie planowania marketingowego	Innowacyjność	177	-0,023	0,766

Źródło: Opracowanie własne.

Wyniki zaprezentowane w tabeli 18 pokazują, że nie występuje istotny związek pomiędzy poziomem zaawansowania planowania marketingowego z wykorzystaniem AI a kluczowymi wskaźnikami efektywności funkcjonowania przedsiębiorstw. Wartości współczynników korelacji Pearsona dla wszystkich analizowanych par zmiennych – satysfakcji klientów, rentowności i innowacyjności – są zbliżone do zera (od -0,006 do -0,025). Oznacza to, że nawet przy relatywnie dużej liczbie obserwacji ($N = 177$) wzrost indeksu AI nie przekłada się statystycznie na wzrost żadnego z analizowanych wyników biznesowych. Dodatkowo, p-wartości ($p = 0,935$ dla satysfakcji klientów, $p = 0,736$ dla rentowności, $p = 0,766$ dla innowacyjności) znacząco przekraczają poziom istotności

0,05, co potwierdza brak zależności między zmiennymi. Wynika z tego, że siła związku między wdrożeniem AI a osiąganymi wynikami nie różni się od efektów przypadkowych.

Z perspektywy teoretycznej i praktycznej wyniki te wskazują, że poziom zaawansowania wykorzystania AI w badanych przedsiębiorstwach nie stanowi czynnika różnicującego ich skuteczność biznesową – ani w obszarze satysfakcji klientów, ani efektywności ekonomicznej, ani innowacyjnego rozwoju. Możliwym wyjaśnieniem braku zależności jest wczesny etap wdrożeń technologii AI w polskich firmach, który nie przynosi jeszcze wymiernych efektów, lub też wpływ AI może być osłabiany przez inne czynniki organizacyjne nieuwzględnione w modelu analizy.

Reasumując, hipoteza pierwsza nie znajduje empirycznego potwierdzenia – poziom zaawansowania planowania marketingowego opartego na AI nie wiąże się ze wzrostem efektów biznesowych w żadnej z analizowanych kategorii.

Hipoteza 2 zakładała, że przedsiębiorstwa z branży technologicznej charakteryzują się wyższym poziomem zaawansowania sztucznej inteligencji niż firmy z innych sektorów gospodarki.

Tabela 19. Poziom zaawansowania AI (indeks ciągły) w różnych branżach

Branża	N	Średnia indeksu AI	Odchylenie std.
Usługi	62	1,91	1,10
Handel	54	1,72	1,03
Technologie	21	2,24	1,19
Inne	12	1,30	1,08
Produkcja	28	1,80	1,08

Wyniki ANOVA: $F(4,172) = 1,73$; $p = 0,145$

Wyniki Kruskal-Wallis: $H = 6,88$; $p = 0,142$

Źródło: Opracowanie własne.

W analizie porównano poziom zaawansowania AI w pięciu głównych branżach reprezentowanych w próbie badawczej. Najwyższy średni poziom indeksu AI odnotowano w firmach technologicznych (2,24), co wskazuje na ich większe zaawansowanie we wdrażaniu rozwiązań sztucznej inteligencji w porównaniu z pozostałymi sektorami (średnie wartości mieszczą się w przedziale od 1,30 do 1,91) (tab. 19).

Pomimo obserwowanej przewagi sektora technologicznego pod względem wartości średnich, ani analiza wariancji ANOVA ($p = 0,145$), ani nieparametryczny test Kruskala-Wallisa ($p = 0,142$) nie potwierdziły istotnych statystycznie różnic między branżami. Oznacza to, że zróżnicowanie poziomu zaawansowania AI pomiędzy analizowanymi sektorami nie przekracza granicy losowej zmienności.

Można zatem stwierdzić, że choć przedsiębiorstwa technologiczne wykazują tendencję do wyższego poziomu wdrożenia AI, różnice te nie są na tyle silne, by uznać je za istotne statystycznie. Hipoteza H2 nie znajduje zatem pełnego potwierdzenia empirycznego.

Uzyskane wyniki implikują, że wdrażanie sztucznej inteligencji w polskich przedsiębiorstwach przebiega w zbliżonym tempie niezależnie od branży. Ewentualna przewaga firm technologicznych może wynikać raczej z ich naturalnych kompetencji w obszarze nowych technologii niż z rzeczywistego efektu branżowego. Można przypuszczać, że w większej próbie lub przy bardziej szczegółowym podziale branż różnice te mogłyby okazać się bardziej wyraźne.

Hipoteza 3 zakładała, że korzystanie z wybranych narzędzi sztucznej inteligencji, takich jak Mailchimp, istotnie wpływa na wyniki marketingowe i rentowność przedsiębiorstw.

W celu weryfikacji hipotezy przeprowadzono analizę wariancji (ANOVA) oraz obliczono współczynniki korelacji punktowo-dwuseryjnej dla trzech kluczowych wskaźników: rentowności, innowacyjności oraz satysfakcji klientów (tab. 20).

Tabela 20. Wyniki biznesowe w zależności od korzystania z Mailchimp

Wynik	N (nie korzysta)	N (korzysta)	Śr. (nie korzysta)	Śr. (korzysta)	SD (nie korzysta)	SD (korzysta)	ANOVA F	ANOVA p	r punktowo-dwuseryjny	p punktowo-dwuseryjny
Rentowność	119	58	2,83	2,34	1,53	1,42	4,13	0,044	-0,15	0,044
Innowacyjność	119	58	3,10	2,50	1,59	1,50	5,77	0,017	-0,18	0,017
Satysfakcja klientów	119	58	3,06	2,62	1,65	1,39	3,03	0,083	-0,13	0,083

Źródło: Opracowanie własne.

Przeprowadzona analiza wykazała istotne statystycznie różnice między firmami korzystającymi z Mailchimp a tymi, które z niego nie korzystają – w dwóch spośród trzech analizowanych wskaźników. Zarówno testy ANOVA, jak i współczynniki korelacji punktowo-dwuseryjnej potwierdziły, że rentowność ($p = 0,044$; $r = -0,15$) oraz innowacyjność ($p = 0,017$; $r = -0,18$) firm używających tego narzędzia są istotnie niższe w porównaniu z pozostałymi przedsiębiorstwami. W przypadku satysfakcji klientów różnice nie osiągnęły poziomu istotności ($p = 0,083$), choć obserwowany kierunek efektu pozostaje spójny z wcześniejszymi wynikami.

Efekty te mają umiarkowaną siłę, jednak przy liczebności próby $N = 177$ uzyskane wartości p świadczą o ich statystycznej istotności. Wskazuje to, że w analizowanej próbie korzystanie z narzędzia Mailchimp wiąże się raczej z niższymi niż wyższymi wskaźnikami rentowności i innowacyjności. Możliwym wyjaśnieniem tego zjawiska jest fakt, że wdrożenie narzędzi AI wymaga odpowiednich kompetencji analitycznych, integracji z procesami marketingowymi oraz dojrzałości organizacyjnej – ich brak może prowadzić do nieskutecznego wykorzystania automatyzacji.

Hipoteza H3 nie znajduje zatem potwierdzenia empirycznego. W badanej próbie korzystanie z narzędzi AI (Mailchimp) nie przekłada się na poprawę wyników biznesowych, a wręcz może je pogarszać. W przyszłych badaniach warto rozszerzyć analizę o inne narzędzia marketingowe, większe próby oraz uwzględnić zmienne kontekstowe, takie jak poziom integracji AI z procesami firmy czy branżowe zróżnicowanie wdrożeń.

Hipoteza 4 zakładała, że wielkość przedsiębiorstwa moderuje wpływ zaawansowania sztucznej inteligencji na wyniki biznesowe. Oznacza to, że efekt wdrożenia AI na rezultaty działalności może być różny w zależności od tego, czy organizacja jest mikro-, małą, średnią czy dużą firmą.

W celu weryfikacji hipotezy przeprowadzono analizę moderacji w modelu regresji liniowej, w którym zmienną objaśnianą stanowiły wyniki biznesowe (rentowność, satysfakcja klientów, innowacyjność), a zmiennymi objaśniającymi — poziom zaawansowania AI, wielkość firmy oraz efekt interakcji pomiędzy nimi. Analiza miała na celu ustalenie czy wpływ zaawansowania AI na wyniki biznesowe różni się istotnie pomiędzy firmami o odmiennych rozmiarach (tab. 21).

Tabela 21. Wartości p dla efektów interakcji (AI × wielkość firmy) w predykcji wyników biznesowych

Wynik	Małe	Mikro	Średnie
Rentowność	p = 0,663	p = 0,952	p = 0,619
Satysfakcja	p = 0,076	p = 0,153	p = 0,709
Innowacyjność	p = 0,064	p = 0,814	p = 0,419

Źródło: Opracowanie własne.

W żadnym z analizowanych modeli nie zaobserwowano istotnego statystycznie efektu moderacji wielkości firmy na związek pomiędzy poziomem zaawansowania AI a wynikami biznesowymi. Wszystkie wartości p przekraczają próg istotności 0,05. Najniższe wartości (dla satysfakcji klientów i innowacyjności w grupie małych firm) zbliżają się do tego progu, jednak nie osiągają poziomu pozwalającego na potwierdzenie efektu.

Wyniki te wskazują, że wpływ wdrożenia AI na rezultaty działalności przedsiębiorstw nie różni się w zależności od ich wielkości. Zarówno mikro-, małe, jak i średnie firmy wykazują podobny schemat zależności — poziom zaawansowania AI nie przekłada się na wyniki biznesowe w odmienny sposób w żadnej z analizowanych grup.

Podsumowując, hipoteza H4 została odrzucona. Analizy nie potwierdziły, że wielkość przedsiębiorstwa w sposób istotny statystycznie moderuje wpływ zaawansowania AI na rentowność, satysfakcję klientów ani innowacyjność w badanej próbie. Wynik ten sugeruje, że czynniki determinujące skuteczność wdrożeń AI są bardziej złożone i mogą zależeć od innych zmiennych, takich jak poziom kompetencji technologicznych, kultura organizacyjna czy stopień integracji AI z procesami zarządczymi.

Hipoteza 5 zakładała, że zaawansowane praktyki marketingowe – takie jak monitorowanie potrzeb klientów w czasie rzeczywistym, alokacja zasobów marketingowych oraz wykorzystanie raportów i analiz opartych na sztucznej inteligencji – są pozytywnie powiązane z wynikami biznesowymi organizacji.

W celu weryfikacji hipotezy przeprowadzono analizę regresji liniowej, w której wymienione praktyki marketingowe uwzględniono jednocześnie jako predyktory trzech kluczowych wyników biznesowych: rentowności, satysfakcji klientów oraz innowacyjności przedsiębiorstw (tab. 22).

Tabela 22. Wyniki regresji wieloczynnikowej (OLS) dla praktyk marketingowych jako predyktorów wyników biznesowych

Wynik biznesowy	Predyktor	β	t	p	95% CI	R ²	Adj. R ²	F (model)	p (model)
Rentowność	Monitorowanie potrzeb klientów w czasie rzeczywistym	0,02	0,21	0,835	[-0,156; 0,193]	0,07	0,05	4,19	0,0068
	Alokacja zasobów marketingowych	0,15	2,05	0,042	[0,006; 0,288]				
	Raporty i analizy AI	0,20	2,11	0,036	[0,013; 0,381]				
Satysfakcja klientów	Monitorowanie potrzeb klientów w czasie rzeczywistym	0,00	0,02	0,985	[-0,182; 0,186]	0,05	0,04	3,28	0,022
	Alokacja zasobów marketingowych	0,18	2,41	0,017	[0,033; 0,330]				
	Raporty i analizy AI	0,15	1,57	0,118	[-0,040; 0,349]				
Innowacyjność	Monitorowanie potrzeb klientów w czasie rzeczywistym	-0,03	-0,30	0,762	[-0,216; 0,158]	0,03	0,01	1,60	0,192
	Alokacja zasobów marketingowych	0,17	2,17	0,031	[0,015; 0,317]				
	Raporty i analizy AI	0,05	0,46	0,649	[-0,152; 0,243]				

Źródło: Opracowanie własne.

W modelu przewidującym rentowność firm uzyskano istotny wynik $F(3,173) = 4,19$; $p = 0,0068$, co oznacza, że zróżnicowanie praktyk marketingowych istotnie tłumaczy część zmienności rentowności. Choć poziom wyjaśnionej wariancji jest umiarkowany ($R^2 = 0,068$; Adj. $R^2 = 0,052$), wskazuje on na znaczenie jakości zarządzania marketingiem w osiąganiu lepszych wyników finansowych. Najważniejszymi predyktorami okazały się: alokacja zasobów marketingowych ($\beta = 0,15$; $p = 0,042$) oraz korzystanie z raportów i analiz AI ($\beta = 0,20$; $p = 0,036$). Monitorowanie potrzeb klientów nie wykazało istotnego wpływu ($p = 0,835$).

Model prognozujący satysfakcję klientów również okazał się istotny ($F(3,173) = 3,28$; $p = 0,022$; $R^2 = 0,054$; $\text{Adj. } R^2 = 0,037$). Najsilniejszym czynnikiem predykcyjnym była ponownie alokacja zasobów marketingowych ($\beta = 0,18$; $p = 0,017$), co sugeruje, że umiejętne gospodarowanie zasobami w obszarze marketingu bezpośrednio przekłada się na wzrost satysfakcji klientów. Pozostałe praktyki – monitorowanie potrzeb oraz raporty AI – nie osiągnęły poziomu istotności ($p = 0,985$ i $p = 0,118$).

W przypadku innowacyjności organizacji model jako całość nie był istotny statystycznie ($F(3,173) = 1,60$; $p = 0,192$; $R^2 = 0,027$; $\text{Adj. } R^2 = 0,01$). Jednakże, w analizie poszczególnych predyktorów znaczący okazał się wpływ alokacji zasobów marketingowych ($\beta = 0,17$; $p = 0,031$), co potwierdza jej znaczenie również w kontekście rozwoju innowacyjności przedsiębiorstw.

Podsumowując, uzyskane wyniki wskazują, że alokacja zasobów marketingowych jest najważniejszą z analizowanych praktyk, mającą pozytywny wpływ na wszystkie trzy wyniki biznesowe – rentowność, satysfakcję klientów i innowacyjność. Dodatkowo, korzystanie z raportów i analiz AI istotnie wspiera wzrost rentowności, co potwierdza znaczenie danych i analityki w efektywnym zarządzaniu marketingiem. Natomiast monitorowanie potrzeb klientów, mimo intuicyjnego znaczenia, nie uzyskało potwierdzenia jako czynnik istotny statystycznie.

Wnioski te potwierdzają, że wdrażanie zaawansowanych praktyk marketingowych sprzyja poprawie wyników biznesowych polskich przedsiębiorstw, szczególnie w wymiarze finansowym i relacyjnym. Choć siła zależności jest umiarkowana, ich konsekwencje praktyczne są istotne i mogą stanowić punkt wyjścia do dalszych badań nad efektywnością integracji analityki AI z procesami marketingowymi.

Dyskusja wyników i implikacje dla teorii oraz praktyki w przedsiębiorstwach SSE

Wyniki przeprowadzonych badań wskazują, że proces wdrażania sztucznej inteligencji w przedsiębiorstwach, w tym w szczególności w organizacjach działających w ramach Specjalnych Stref Ekonomicznych, ma charakter złożony i przebiega w sposób stopniowy. Zebrane dane empiryczne potwierdzają, że większość badanych podmiotów znajduje się na wczesnym etapie transformacji cyfrowej, a wykorzystanie rozwiązań AI koncentruje się głównie na obszarach o niskim poziomie zaawansowania, takich jak automatyzacja procesów, analiza danych czy wstępne wsparcie działań marketingowych.

Jednym z najważniejszych wniosków płynących z analizy jest fakt, że przedsiębiorstwa z sektora SSE wykazują znaczące zróżnicowanie pod względem poziomu zaawansowania wdrożeń AI. Choć część z nich deklaruje wykorzystanie rozwiązań sztucznej inteligencji w codziennej działalności operacyjnej, większość respondentów ocenia stopień integracji technologii z kluczowymi procesami biznesowymi jako raczej niski lub średni. Szczególnie widoczne jest to w obszarach wymagających zaawansowanego zarządzania danymi, bezpieczeństwa informacji czy integracji AI z planowaniem strategicznym.

Dane empiryczne wskazują również, że wśród przedsiębiorstw działających w SSE brakuje pełnego powiązania pomiędzy wdrażaniem AI a strategią organizacyjną. W wielu przypadkach działania te mają charakter incydentalny lub eksperymentalny, a nie strategiczny i długofalowy. Tylko nieliczne organizacje zadeklarowały, że sztuczna inteligencja stanowi istotny element ich wizji i planów rozwoju. Świadczy to o potrzebie większego zaangażowania kadry zarządzającej w procesy cyfrowej transformacji oraz budowania świadomości na temat potencjału AI jako narzędzia wspierającego wzrost efektywności i konkurencyjności przedsiębiorstw.

Z punktu widzenia praktyki gospodarczej, uzyskane wyniki sugerują konieczność wzmocnienia działań w zakresie rozwoju kompetencji cyfrowych i analitycznych wśród pracowników oraz menedżerów przedsiębiorstw. Skuteczne wdrażanie AI wymaga nie tylko odpowiedniej infrastruktury technologicznej, ale przede wszystkim przygotowania organizacyjnego, obejmującego kulturę otwartą na innowacje, gotowość do eksperymentowania oraz umiejętność adaptacji do zmieniających się warunków rynkowych. W tym kontekście szczególne znaczenie ma również współpraca międzysektorowa – łączenie potencjału przedsiębiorstw, instytucji naukowych i administracji publicznej w celu transferu wiedzy, komercjalizacji technologii oraz zwiększenia efektywności wykorzystania dostępnych zasobów.

Z teoretycznego punktu widzenia, uzyskane wyniki potwierdzają aktualność koncepcji zdolności dynamicznych, zgodnie z którą przewaga konkurencyjna organizacji wynika z jej umiejętności nieustannego uczenia się, adaptacji i rekonfiguracji posiadanych zasobów w reakcji na zmiany otoczenia. W przypadku przedsiębiorstw z SSE, zdolność ta jest szczególnie istotna, ponieważ otoczenie tych organizacji charakteryzuje się dużą dynamiką, konkurencyjnością i presją technologiczną.

Wnioski z badań mają również znaczenie praktyczne dla polityki rozwoju Specjalnych Stref Ekonomicznych. Wynika z nich potrzeba dalszego wspierania

procesów cyfrowej transformacji poprzez programy szkoleniowe, doradcze i inwestycyjne, które umożliwią przedsiębiorstwom skuteczniejsze wdrażanie technologii AI. SSE mogą odgrywać rolę **lokalnych centrów innowacji**, wspierających przedsiębiorców w dostępie do technologii, wiedzy eksperckiej oraz rozwiązań infrastrukturalnych sprzyjających rozwojowi inteligentnych modeli biznesowych.

Podsumowując, przeprowadzone badania empiryczne potwierdzają, że sztuczna inteligencja staje się coraz istotniejszym elementem funkcjonowania przedsiębiorstw, jednak jej wdrażanie w wielu przypadkach pozostaje ograniczone do podstawowych zastosowań operacyjnych. Dalszy rozwój w tym obszarze wymaga nie tylko inwestycji technologicznych, ale także kształtowania odpowiednich postaw organizacyjnych i kompetencji menedżerskich, które umożliwią pełne wykorzystanie potencjału AI w środowisku gospodarczym Specjalnych Stref Ekonomicznych.

Zakończenie

Celem niniejszej pracy była kompleksowa analiza gotowości technologicznej i stopnia wykorzystania sztucznej inteligencji w przedsiębiorstwach, ze szczególnym uwzględnieniem specyficznego kontekstu, jakim są Specjalne Strefy Ekonomiczne w Polsce. Podjęto próbę zdiagnozowania kluczowych czynników – infrastrukturalnych, strategicznych i aplikacyjnych – które determinują sukces wdrożeń AI, a także empirycznej weryfikacji ich związku z osiąganymi wynikami biznesowymi. Realizacja tak sformułowanego celu wymagała zarówno pogłębionych studiów literaturowych, jak i przeprowadzenia badań empirycznych.

W części teoretycznej pracy wykazano, że skuteczna transformacja w kierunku organizacji opartej na danych jest procesem holistycznym, opartym na trzech nierozdzielnie połączonych filarach. Pierwszym jest dojrzałość infrastrukturalna, rozumiana nie jako suma odizolowanych technologii, lecz jako synergiczny ekosystem obejmujący nowoczesne architektury danych (np. Data Lakehouse), skalowalne zasoby obliczeniowe (chmura, GPU/TPU) oraz solidne ramy bezpieczeństwa. Drugim filarem jest dojrzałość strategiczna, która wykracza poza samą technologię i obejmuje integrację AI z celami biznesowymi, zaangażowanie przywództwa (AI Leadership), zdolność do pokonywania barier kulturowych oraz promowanie kultury eksperymentowania. Trzeci filar to zdolność do aplikacji AI w kluczowych procesach biznesowych, co zilustrowano na przykładzie marketingu, gdzie AI rewolucjonizuje każdy etap zarządzania informacją, prowadząc do hiperpersonalizacji i redefinicji doświadczenia klienta.

Nadrzędny wniosek płynący z przeprowadzonych analiz jest jednoznaczny: sukces w erze sztucznej inteligencji nie jest pochodną samego posiadania zaawansowanej technologii, lecz wynikiem organizacyjnej dojrzałości do jej strategicznego wykorzystania. Zidentyfikowana w Polsce "luka wdrożeniowa" nie jest luką technologiczną, lecz przede wszystkim luką kompetencyjną, strategiczną i kulturową. To zdolność do holistycznego zarządzania zmianą, integracji AI z modelem

biznesowym i budowania kultury opartej na danych decyduje o tym, czy potężne narzędzia analityczne stają się motorem wzrostu, czy jedynie kosztownym, nie-wykorzystanym zasobem. Kontekst Specjalnych Stref Ekonomicznych w dobitny sposób unaoczniał ten paradoks – samo stworzenie sprzyjającego "hardware'u" (infrastruktury, ulg finansowych) jest niewystarczające bez inwestycji w "software" (kompetencje, przywództwo, kultura innowacji).

Należy jednak podkreślić, że niniejsza praca, jak każde badanie naukowe, posiada pewne ograniczenia. Po pierwsze, ma ona charakter przekrojowy, co oznacza, że analizuje zjawiska w jednym, konkretnym punkcie czasowym. Pozwala to na identyfikację korelacji, lecz nie na jednoznaczne wnioskowanie o związkach przyczynowo-skutkowych, które wymagałoby badań o charakterze podłużnym. Po drugie, badanie opierało się na metodzie ankietowej, co wiąże się z ryzykiem subiektywizmu ocen respondentów. Po trzecie, próba badawcza, choć starannie dobrana, jest ograniczona do określonej liczby przedsiębiorstw, co nakazuje ostrożność w generalizowaniu wniosków na całą populację firm w Polsce.

Zidentyfikowane ograniczenia, a także niezwykle dynamika samego zjawiska, otwierają szerokie pole dla przyszłych badań. Wskazane byłoby przeprowadzenie badań longitudinalnych, śledzących ewolucję dojrzałości AI w firmach na przestrzeni kilku lat. Cenne poznawczo byłyby również pogłębione studia przypadku (case studies), które pozwoliłyby na jakościowe zrozumienie mechanizmów stojących za sukcesem lub porażką transformacji cyfrowej. Dalsze badania mogłyby również objąć analizę porównawczą różnych sektorów gospodarki oraz zbadać wpływ AI na inne obszary funkcjonalne, takie jak finanse, logistyka czy zarządzanie zasobami ludzkimi.

Wyniki pracy niosą za sobą również istotne implikacje praktyczne. Dla menedżerów głównym przesłaniem jest konieczność myślenia o AI nie jako o projekcie IT, lecz jako o strategicznej transformacji całego przedsiębiorstwa, wymagającej inwestycji we wszystkie trzy filary: infrastrukturę, strategię i ludzi. Dla decydentów politycznych, w szczególności zarządzających Specjalnymi Strefami Ekonomicznymi, płynie stąd wniosek, że instrumenty wsparcia powinny ewoluować, przenosząc punkt ciężkości z zachęt finansowych na programy budujące konkretne kompetencje cyfrowe, wspierające transfer wiedzy i promujące kulturę innowacji.

Podsumowując, sztuczna inteligencja bez wątpienia stanowi jedną z najbardziej transformacyjnych sił we współczesnej gospodarce. Jednak jej potencjał zostanie w pełni uwolniony nie przez te organizacje, które najszybciej zakupią nową technologię, lecz przez te, które najskuteczniej nauczą się myśleć strategicznie,

ZAKOŃCZENIE

działać odważnie i adaptować się w sposób ciągły. Zbudowanie takiej zdolności jest dziś największym wyzwaniem i jednocześnie największą szansą dla polskich przedsiębiorstw w dążeniu do globalnej konkurencyjności.

BIBLIOGRAFIA

- Adesina, A. A., Iyelolu, T. V., & Paul, P. O. (2024). Leveraging predictive analytics for strategic decision-making: Enhancing business performance through data-driven insights. *World Journal of Advanced Research and Reviews*, 22(3), 1927–1934.
- Adwani, A. (2025). Predictive Analytics for Business Strategy: Leveraging Machine Learning for Competitive Advantage. *International Journal of Social Impact*, 10(3), 295-307.
- Agbaakin, Q. A. (2025). Leveraging Artificial Intelligence as a Strategic Growth Catalyst for Small and Medium-sized Enterprises. *arXiv*. <https://arxiv.org/html/2509.14532v1>.
- Ahuja, A. (2024). *Pioneering enterprise architecture: Transforming global enterprises. Strategies for digital transformation & scalability*. New York, NY: Ashutosh Ahuja.
- Ahuja, N. (2024). AI governance and ethical considerations in the age of generative models. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5007984.
- Akter, J., Roy, A., Rahman, S., Mohona, S., & Ara, J. (2025). Artificial intelligence-driven customer lifetime value (CLV) forecasting: Integrating RFM analysis with machine learning for strategic customer retention. *Journal of Computer Science and Technology Studies*, 7(1), 249-257.
- Akter, S., Sultana, S., Mariani, M., Wamba, S. F., Spanaki, K., & Dwivedi, Y. K. (2023). Advancing algorithmic bias management capabilities in AI-driven marketing analytics research. *Industrial Marketing Management*, 114, 243–261.

- Al Khaldy, M. A., Al-Obaydi, B. A. A., & al Shari, A. J. (2023). The impact of predictive analytics and AI on digital marketing strategy and ROI. Posted November 22, 2023. https://www.researchgate.net/publication/375423745_The_Impact_of_Predictive_Analytics_and_AI_on_Digital_Marketing_Strategy_and_ROI.
- Alevizos, L., & Dekker, M. (2024). Towards an AI-enhanced cyber threat intelligence processing pipeline. *Electronics*, 13(11), 2021.
- Alhasan, M. (2025). Predictive analytics in marketing: Evaluating its effectiveness in driving customer engagement. *Advances in Science, Technology and Engineering Systems Journal*, 10(3), 45–51.
- Ali, A., & Dornaika, F. (2025). Edge Artificial Intelligence: A Systematic Review of Evolution, Taxonomic Frameworks, and Future Horizons. <https://arxiv.org/html/-2510.01439v1>.
- Amola, P. (2025). AI and data lakes: Transforming enterprise data storage and retrieval with machine learning. https://www.researchgate.net/publication/389433288_AI_and-_Data_Lakes_Transforming_Enterprise_Data_Storage_and_Retrieval_with_Machine_Learning.
- Armbrust, M., Ghodsi, A., Xin, R., & Zaharia, M. (2021). Lakehouse: A new generation of open platforms that unify data warehousing and advanced analytics. https://www.cidrdb.org/cidr2021/papers/cidr2021_paper17.pdf.
- Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Benetton, A., Tabik, S., & Barbado, A. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115.
- Ashfaq, B. (2025). Why GPUs Rule AI and Graphics: The Ultimate Guide to Parallel Computing. <https://medium.com/@ashfaqbs/why-gpus-rule-ai-and-graphics-the-ultimate-guide-to-parallel-computing-f6109ae166de>.
- Astrup, N. (2020). National strategy for artificial intelligence. Norwegian Ministry of Local Government and Modernisation. https://www.regjeringen.no/contentassets/-1febbbb2c4fd4b7d92c67ddd353b6ae8/en-gb/pdfs/ki-strategi_en.pdf.
- Aws.amazon.com (2024). What are transformers in artificial intelligence. Retrieved from <https://aws.amazon.com/what-is/transformers-in-artificial-intelligence/>.

- Azizi, A., Mohammadi, M. Q., & Samadzai, A. W. (2024). AI in cyber defense: Privacy risks, public trust, and policy challenges. *Jurnal Ilmiah Dinamika Sosial*, 9(1), 103–118.
- Balasubramaniam, N., Kauppinen, M., Rannisto, A., Hiekkanen, K., & Kujala, S. (2023). Transparency and explainability of AI systems: From ethical guidelines to requirements. *Information and Software Technology*, 159, 107197.
- Basal, M., Saraç, E., & Özer, K. (2024). Dynamic pricing strategies using artificial intelligence algorithm. *Open Journal of Applied Sciences*, 14(8), 1963–1978.
- Beem, V. R. (2025). AI-driven personalization in retail: Transforming customer experience through intelligent product recommendations. <https://doi.org/10.37745/-ejcsit.2013/vol13n38117131>.
- Belcastro, L., Marozzo, F., Orsino, A., Talia, D., & Trunfio, P. (2025). Navigating the Edge-Cloud Continuum: A State-of-Practice Survey. <https://arxiv.org/html/2506.02003v1>.
- Bell, C., Olukemi, A., & Brooklyn, P. (2024). AI-Driven personalization in digital marketing: Effectiveness and ethical considerations. SSRN Pre-print. <https://doi.org/10.-2139/ssrn.4906214>.
- Bergman, D., Huang, T., Brooks, P., Lodi, A., & Raghunathan, A. U. (2019). JANOS: An integrated predictive and prescriptive modeling framework. *arXiv*. <https://arxiv.org/abs/1911.09461>.
- Bevilacqua, M., Berente, N., Domin, J., Goehring, J., & Rossi, A. (2023). The Return on Investment in AI Ethics: A Holistic Framework. *arXiv*. <https://arxiv.org/pdf/2309.13057>.
- Bevilacqua, M., Masárová, A., Perotti, L., & Ferraris, A. (2025). Enhancing top managers' leadership with artificial intelligence: Insights from a systematic literature review. *Review of Managerial Science*, 19, 2899–2935.
- Bomma, H. (2024). AI Integrated Data Governance and DataLineage. <https://www.ijssat.org/papers/2024/2/1645.pdf>.
- Boozary, P., Sheykhan, S., GhorbanTanhaei, H., & Magazzino, C. (2025). Enhancing customer retention with machine learning: A comparative analysis of ensemble models for accurate churn prediction. *International Journal of Information Management Data Insights*, 5(1), 100331.

- Boteju, S., Banbaduge, C., Vatsalan, D., & Arachchilage, N. A. G. (2023). SoK: Demystifying Privacy Enhancing Technologies Through the Lens of Software Developers. <https://arxiv.org/pdf/2401.00879>.
- Brewis, C., Dibb, S., & Meadows, M. (2023). Leveraging big data for strategic marketing: A dynamic capabilities model for incumbent firms. *Technological Forecasting and Social Change*, 190, 122402.
- Brey, P., & Dainow, B. (2024). Ethics by design for artificial intelligence. *AI and Ethics*, 4(4), 1265–1277.
- ByteBridge (2025). GPU and TPU comparative analysis report. <https://byte-bridge.-medium.com/gpu-and-tpu-comparative-analysis-report-a5268e4f-0d2a>.
- Calvi, A., Malgieri, G., & Kotzinos, D. (2024). The unfair side of Privacy Enhancing Technologies: addressing the trade-offs between PETs and fairness. <https://facctconference.org/static/papers24/facct24-139.pdf>.
- Carmichael, M. (2025). Demonstrating AI ROI: How to Measure and Prove the Value of Your AI Investments. ISACA. <https://www.isaca.org/resources/news-and-trends/newsletters/atisaca/2025/volume-5/how-to-measure-and-prove-the-value-of-your-ai-investments>.
- Celi, L., & Miles, R. (2020). Driving ROI through AI. Econsult Solutions Inc. https://econsultsolutions.com/wp-content/uploads/2020/09/ESITL_Driving-ROI-through-AI_FINAL_September-2020.pdf.
- Chen, K. (2025). What Is Cloud Economics? The Ultimate Guide. <https://www.-oracle.com/cloud/cloud-economics-explained/>.
- Chinthapatla, A. (2023). AWS SageMaker vs Google Cloud AI: Unveiling the powerhouses of machine learning. https://www.researchgate.net/publication/379262637-_AWS_SageMaker_vs_Google_Cloud_AI_Unveiling_the_Powerhouses_of_Machine_Learning.
- Chojnowski, M. (2025). Godna zaufania AI. Jak używać sztucznej inteligencji zgodnie z wytycznymi Komisji Europejskiej w zakresie etyki?. https://ai.gov.pl/media/2025/05/-Raport_GRAI_godna_zaufania_sztuczna_inteligencja_styczen_2025_pdf.pdf.
- Cichy, A., Czapłuk, D., Kiebzak, P., & Kieroński, P. (2024). Rynek pracy a rewolucja cyfrowa. Wpływ AI i nowych technologii. https://www.politykainsight.pl/_resource/-multimedia/20360905.

- Climent, R. C., Haftor, D. M., & Staniewski, M. W. (2024). AI-enabled business models for competitive advantage. *Journal of Innovation & Knowledge*, 9, 100532.
- Costa, C. J., Aparicio, M., Aparicio, S., & Aparicio, J. T. (2024). The democratization of artificial intelligence: Theoretical framework (s. 9). *Applied Sciences*, 14(18), 8236.
- Cowger, J. (2023). AI governance and corporate responsibility: Examining emerging frameworks. *Journal of Law, Technology & the Internet*, 14(2), 145.
- Damnjanović, A. M., Rašković, M., Janković, S. D., Jevtić, B., Skoropad, V. N., Marković, Z. D., Lukić-Vujadinović, V., Injac, Z., & Marinković, S. (2025). AI-enabled strategic transformation and sustainable outcomes in Serbian SMEs. *Sustainability*, 17(19), 8672.
- Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2019). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42.
- Dąbrowski, M. (2023). Zastosowanie architektury Nvidia CUDA do przyspieszania procesu uczenia sieci neuronowych. <https://repo.pw.edu.pl/docstore/download.seam?-entityType=bachelor&entityId=WUT307412&fileId=WUT307411>
- Dilmegani, C. (2025). Data quality in AI: Challenges and best practices. <https://research.aimultiple.com/data-quality-ai/>.
- Ding, Z., Fu, Q., Ding, J., Deng, G., Liu, Y., & Li, Y. (2025). A rusty link in the AI supply chain: Detecting evil configurations in model repositories. <https://arxiv.org/-html/2505.01067v1>.
- Dorgbefe, E. A. (2021). Enhancing customer retention using predictive analytics and personalization in digital marketing campaigns. *International Journal of Science and Research Archive*, 4(1), 403–423.
- Douaioui, K., Oucheikh, R., Benmoussa, O., & Mabrouki, C. (2024). Machine learning and deep learning models for demand forecasting in supply chain management: A critical review. *Applied System Innovation*, 7(5), 93.
- Dubey, R. (2025). The data lakehouse: An evolving paradigm in data architecture. https://www.researchgate.net/publication/393759319_The_Data_Lakehouse_An_Evolving_Paradigm_in_Data_Architecture.

- Dzreke, S. S. (2025). The competitive advantage of AI in business: A strategic imperative (pp. 16). *International Journal for Multidisciplinary Research*, 7(4).
- Eisenberg, J., Gamboa, J., & Sherman, J. (2025). The Unified Control Framework: Establishing a Common Foundation for Enterprise AI Governance, Risk Management and Regulatory Compliance. *arXiv*. <https://arxiv.org/pdf/2503.05937>.
- Elad, B. (2025). AI in Marketing Statistics 2025: ROI, Tools & Trends. *SQ Magazine*. <https://sqmagazine.co.uk/ai-in-marketing-statistics/>.
- Exacto. (2024). Jak wykorzystać AI w komunikacji z klientem? <https://www.exacto.pl/jak-wykorzystac-ai-w-komunikacji-z-klientem/>.
- EY Polska. (2023). Customer experience, czyli zarządzanie doświadczeniem klienta. https://www.ey.com/pl_pl/insights/digital-first/doswiadczenie-klienta-customer-experience-czyli-zarzadzanie-doswiadczeniem-klienta.
- Fenwick, M., Vermeulen, E. P. M., & Compagnucci, L. (2024). AI governance and the “siren call” of human oversight. *arXiv*. <https://arxiv.org/abs/2407.19439>.
- Frąckiewicz, M. (2025). Wyścig do 6G: Jak sieć nowej generacji zrewolucjonizuje łączność i pozostawi 5G daleko w tyle. <https://ts2.tech/pl/wyscig-do-6g-jak-siec-nowej-generacji-zrewolucjonizuje-lacznosc-i-pozostawi-5g-daleko-w-tyle/>.
- Fu, M., Pasuksmit, J., & Tantithamthavorn, C. (2024). AI for DevSecOps: A landscape and future opportunities. <https://arxiv.org/abs/2404.04839>.
- GhorbanTanhaei, H., Boozary, P., Sheykhan, S., Rabiee, M., Rahmani, F., & Hosseini, I. (2024). Predictive analytics in customer behavior: Anticipating trends and preferences. *Results in Control and Optimization*, 17, 100462.
- Gonçalves, R., Dias, Á., Lopes da Costa, R., & Pereira, L. (2022). Gaining competitive advantage through artificial intelligence adoption. *International Journal of Electronic Business*, 17(4), 386-406.
- Gontarz, A. (2025). Sztuczna inteligencja w zarządzaniu danymi. <https://crn.pl/artykuly/-sztuczna-inteligencja-w-zarzadzaniu-danymi/>.
- Goswami, M. J. (2021). Challenges and solutions in integrating AI with multi-cloud architectures. *International Journal of Enhanced Research in Management & Computer Applications*, 10(10), 67-73.

- Gregory, B., & Austin, M. (2023). AI and the role of the board of directors. Harvard Law School Forum on Corporate Governance. <https://corpgov.law.harvard.edu/2023/10/07/ai-and-the-role-of-the-board-of-directors/>.
- Gröger, C., Schwarz, H., & Mitschang, B. (2014). Prescriptive Analytics for Recommendation-Based Business Process Optimization. In W. Abramowicz & A. Kokkinaki (Eds.), *Business Information Systems. BIS 2014. Lecture Notes in Business Information Processing, LNBIP*, 176, 25-37.
- Groszkowski, D. (2023). Materiały pokonferencyjne autorstwa uczestników webinaru: „PROJEKTOWANIE SYSTEMÓW SI ZGODNYCH Z RODO”. Urząd Ochrony Danych Osobowych. <https://uodo.gov.pl/pl/file/4380>.
- Grupa Robocza ds. AI przy Kancelarii Prezesa Rady Ministrów (2023). RAPORT. Analiza związku Aktu w sprawie sztucznej inteligencji z wybranymi obowiązującymi i projektowanymi regulacjami prawnymi. <https://www.gov.pl/attachment/6c940b59-4c78-44f3-adfa-8d7fa0bcfd1c>.
- Guo, J., Wang, L., & Yang, X. (2023). Normal Workflow and Key Strategies for Data Cleaning Toward Real-World Data: Viewpoint. <https://pmc.ncbi.nlm.nih.gov/articles/-PMC10557005/>.
- Hahn, D. (2025). Data ingest storage – the unsung hero in AI and LLM pipelines. <https://omdia.tech.informa.com/blogs/2025/sep/data-ingest-storage-the-unsung-hero-in-ai-and-llm-pipelines>.
- Hale, M. (2025). Generative AI reshapes the future of marketing: Key trends, stats, and what’s next for marketers in 2025 and beyond. GSDC. <https://www.gsdcouncil.org/-blogs/generative-ai-future-of-marketing-trends>.
- Helmick, L. (2022). General Purpose Computing on Graphics Processing Units for Accelerated Deep Learning in Neural Networks. <https://digitalcommons.liberty.edu/-cgi/viewcontent.cgi?article=2258&context=honors>.
- Hochkamp, A., Scheidler, A., & Rabe, M. (2025). Review of Maturity Models for Data Mining and Proposal of a Data Preparation Maturity Model Prototype for Data Mining. *Computers*, 14(4), 146.
- Hu, C., Din, Q. M. U., & Tahir, A. (2025). Artificial intelligence symbolic leadership in small and medium-sized enterprises: Enhancing employee flexibility and technology adoption. *Systems*, 13(4), 216.

- Huang, H., Wang, F., Song, M., Baležentis, T., & Streimikiene, D. (2021). Green innovations for sustainable development of China: Analysis based on the nested spatial panel models. *Technology in Society*, 65, 101593.
- Huang, M.-H., & Rust, R. T. (2020). Engaged to a robot? The role of AI in service. *Journal of Service Research*, 24(1), 30-41.
- Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49, 30-50.
- Hunter, D. (2025). The AI marketing revolution: How personalization, automation, and predictive analytics will redefine marketing by 2030. *AIJ Guest Post*. <https://www.aijourn.com/the-ai-marketing-revolution-how-personalization-automation-and-predictive-analytics-will-redefine-marketing-by-2030/>.
- Imani, M., Joudaki, M., Beikmohammadi, A., & Arabnia, H. R. (2025). Customer churn prediction: A systematic review of recent advances, trends, and challenges in machine learning and deep learning. *Machine Learning and Knowledge Extraction*, 7(3), 105.
- Industrial Internet Consortium (2019). The Edge Computing Advantage. https://www.iiconsortium.org/pdf/IIC_Edge_Computing_Advantages_White_Paper_2019-10-24.pdf.
- Jahid, M. (2025). AI-driven optimization and risk modeling in strategic economic zone development for mid-sized economies: A review approach. *ResearchGate*. https://www.researchgate.net/publication/394602488_AI-DRIVEN_OPTIMIZATION-_AND_RISK_MODELING_IN_STRATEGIC_ECONOMIC_ZONE_DEVELOPMENT_FOR_MID-SIZED_ECONOMIES_A_REVIEW_APPROACH.
- Jamarani, E., Haddadi, M., Sarvzadeh, A., Kashani, M. H., Akbari, R., & Moradi, H. (2024). Big data and predictive analytics: A systematic review of applications. *Artificial Intelligence Review*, 57, Article 176.
- Janssen, M. (2022). The Evolution of Data Storage Architectures: Examining the Value of the Data Lakehouse. <https://essay.utwente.nl/essays/92801>.
- Jeffrey, N., Tan, Q., & Villar, J. R. (2023). A review of anomaly detection strategies to detect threats to cyber-physical systems. *Electronics*, 12(15), 3283.

- Jefryy, M., Arman, S., & Takri, R. (2025). AI in cybersecurity: A double-edged sword case. <https://www.techrxiv.org/users/692320/articles/1317665/master/file/data/AI%20-in%20Cybersecurity%20A%20Double-Edged%20Sword-Case/AI%20in%20Cyber-security%20A%20Double-Edged%20Sword-Case.pdf?inline=true>.
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(2).
- John, R. (2025). Cloud-edge AI orchestration: Distributed intelligence models for real-time enterprise decision-making. https://www.researchgate.net/publication/390760320-Cloud-Edge_AI_Orchestration_Distributed_Intelligence_Models_for_Real-Time_Enterprise_Decision-Making.
- Jonker, J., & Gomstyn, A. (2024). What is a data lakehouse? <https://www.ibm.com/-think/topics/data-lakehouse>.
- Jonker, M., & Rogers, A. (2024). What is algorithmic bias? IBM Think Blog. <https://www.ibm.com/think/topics/algorithmic-bias>.
- Jorzik, N., Klein, M., Kanbach, D., & Kraus, S. (2024). AI-driven business model innovation: A systematic review and research agenda. ResearchGate. https://www.-researchgate.net/publication/381482374_AI-driven_business_model_innovation_A_systematic_review_and_research_agenda.
- Jorzik, P., Klein, S. P., Kanbach, D. K., & Kraus, S. (2024). AI-driven business model innovation: A systematic review and research agenda. *Journal of Business Research*, 182, 114764.
- Jouppi, N. P., Young, C., Patil, N., & Patterson, D. (2017). In-Datcenter Performance Analysis of a Tensor Processing Unit. <https://arxiv.org/abs/1704.04760>.
- Jurgens, S., & Dal Cin, P. (2025). Artificial intelligence: The impact of AI on education for all learners. <https://ciddl.org/wp-content/uploads/2025/03/Artificial-Intelligence-The-Impact-of-AI-on-Education-for-All-Learners.pdf>.
- Jurgens, S., & Dal Cin, P. (2025). Global cybersecurity outlook 2025. World Economic Forum. https://reports.weforum.org/docs/WEF_Global_Cybersecurity_Outlook_2025.pdf.
- Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., & Bhagoji, A. N. (2021). Advances and Open Problems in Federated Learning. <https://arxiv.org/abs/1912.04977>.

- Kapuściński, P. (2025). AI w B2B: Jak sztuczna inteligencja zmienia marketing i sprzedaż. TTMS. <https://ttms.com/pl/ai-w-b2b-jak-sztuczna-inteligencja-zmienia-marketing-i-sprzedaz/>.
- Kasinadhsarma, S. (2024). Revolutionizing Machine Learning: The Emergence and Impact of Google's TPU Technology. <https://easychair.org/publications/preprint/-KFMC/open>.
- Katragadda, S. (2025). Barriers to AI Adoption in Strategic IT Planning: Organizational and Technical Perspectives. ResearchGate. https://www.researchgate.net/publication/-393795320_Barriers_to_AI_Adoption_in_Strategic_IT_Planning_Organizational_and_Technical_Perspectives.
- Khadka, D., & Ullah, A. (2025). Human factors in cybersecurity: An interdisciplinary review and framework proposal. *International Journal of Information Security*, 24, 119.
- Khairullah, H., Harris, M., Hadi, F., Sandhu, A., Ahmad, S., & Alshara, M. (2025). Implementing artificial intelligence in academic and administrative processes through responsible strategic leadership in the higher education institutions, 10, 1548104.
- Khamoushi, E. (2024). AI in food marketing from personalized recommendations to predictive analytics: Comparing traditional advertising techniques with AI-driven strategies. arXiv. <https://arxiv.org/abs/2410.01815>.
- Kliś, T. (2025). Marketing oparty na danych. Skuteczność i percepcja algorytmów predykcyjnych wśród konsumentów. *Zarządzanie Mediami*, 13(2), 97–115.
- Kokkonen, T., Lovén, L., Motlagh, N. H., Kumar, T., & Partala, J. (2022). Autonomy and Intelligence in the Computing Continuum: Challenges, Enablers, and Future Directions for Orchestration. <https://arxiv.org/abs/2205.01423>.
- Kong, X., Chen, Z., Liu, W., Ning, K., Zhang, L., Marier, S. M., Liu, Y., Chen, Y., & Xia, F. (2024). Deep learning for time series forecasting: A survey. *International Journal of Machine Learning and Cybernetics*, 16(7-8), 5079–5112.
- Kostro, M. (2021). Chmura obliczeniowa – królowa transformacji technologicznej. <https://mitsmr.pl/transformatcja-organizacyjna/chmura-obliczeniowa-krolowa-transformacji-technologicznej/>.

- Kowalski, A. M., Lewandowska, M. S., & Weresa, M. A. (2023). Polska: Raport o konkurencyjności 2023. Znaczenie przedsiębiorczości w kształtowaniu przewag konkurencyjnych. Oficyna Wydawnicza SGH.
- Krajnc, A., Spielvogel, C., Grahovac, J., Ecsedi, S., Rasul, S., Poetsch, N., & Traub-Weidinger, T. (2022). Automated data preparation for in vivo tumor characterization with machine learning. <https://www.frontiersin.org/journals/oncology/articles/10.3389/fonc.-2022.1017911/full>.
- Krawiec, M. (2025). Od kosztów po kompetencje: Co blokuje rozwój AI w firmach. *MobileTrends*. <https://mobiletrends.pl/od-kosztow-po-kompetencje-co-blokuje-rozwoj-ai-w-firmach/>.
- Krawiec, P. (2025). Korzyści z wdrożenia AI w biznesie: Efektywność, decyzje i nowe źródła przychodów. <https://mobiletrends.pl/korzysci-z-wdrozenia-ai-w-biznesie-efektywnosc-decyzje-i-nowe-zrodla-przychodow/>.
- Kujore, V., Adebayo, A., Sambakui, O., & Segbenu, B. S. (2025). Transformative role of generative AI in marketing content creation and brand-engagement strategies. *GSC Advanced Research & Review*, 23(3), 1–11.
- Kumar, N. (2025). Intelligent customer segmentation: Unveiling consumer patterns with machine learning. *Journal of Umm Al-Qura University for Engineering and Architecture*, 16, 774–783.
- Kumar, V., Ashraf, A. R., & Nadeem, W. (2024). AI-powered marketing: What, where, and how? *International Journal of Information Management*, 77, 102783.
- Lambrecht, A., & Tucker, C. E. (2018). Algorithmic bias? An empirical study into apparent gender-based discrimination in the display of STEM career ads. SSRN. <https://doi.org/10.2139/ssrn.2852260>.
- Lepenioti, K., Bousdekis, A., Apostolou, D., & Mentzas, G. (2020). Prescriptive analytics: Literature review and research challenges. *International Journal of Information Management*, 50, 57-70.
- Liang, L. Y., & Yubin, Z. (2025). The impacts of digital transformation on firm performance: A case study of Tesla business model. *International Journal of Research and Scientific Innovation (IJRSI)*, 12(5), 287-302. <https://www.rsisinternational.-org/journals/ijrsi/articles/the-impacts-of-digital-transformation-on-firm-performance-a-case-study-of-tesla-business-model/>.

- Licea, G. (2025). Business and Regulatory Responses to Artificial Intelligence: Dynamic Regulation, Innovation Ecosystems and the Strategic Management of Disruptive Technology. ResearchGate. https://www.researchgate.net/publication/393804905_AI-Driven_KPIs_for_Measuring_and_Improving_Business-IT_Alignment.
- Liu, J., Zhang, C., Liu, Y., & Zhang, D. (2022). Can the Special Economic Zones Promote the Green Technology Innovation of Enterprises? An Evidence From China. *Frontiers in Environmental Science*, 10, 870019.
- Liu, X., Zhang, Y., Wang, D., Deng, Z., & Wu, B. (2019). Dynamic Pricing on E-commerce Platform with Deep Reinforcement Learning: A Field Experiment. *arXiv*. <https://arxiv.org/abs/1912.02572>.
- Liu-Thompkins, Y., Okazaki, S., & Li, H. (2022). Artificial empathy in marketing interactions: Bridging the human-AI gap in affective and social customer experience. *Journal of the Academy of Marketing Science*, 50, 1198–1218.
- Locascio, D. (2023). Artificial intelligence risk management framework (AI RMF 1.0). National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs-/ai/nist.ai.100-1.pdf>.
- Longo, L., & Saad, R. (2025). Efficient Data Ingestion in Cloud-based architecture: a Data Engineering Design Pattern Proposal. <https://arxiv.org/abs/2503.16079>.
- Lopez-Miguel, A. (2021). Survey on Preprocessing Techniques for Big Data Projects, 7(1), 14.
- Luitse, S. (2024). Platform power in AI: The evolution of cloud infrastructures. <https://policyreview.info/articles/analysis/platform-power-ai-evolution-cloud-infrastructures>.
- Łochowska, A. (2025). Statystyki cyberbezpieczeństwa – dane z 2025 roku. <https://www.exorigo-upos.pl/blog/statystyki-cyberbezpieczenstwa-dane-z-2025-roku/>.
- Maheswaran, K. (2025). GPU Performance Optimization for Deep Learning. <https://www.digitalocean.com/community/tutorials/an-introduction-to-gpu-optimization>.

- Mahu, A., Alamu, O., & Frank, R. (2025). From human to hybrid leadership competencies in an AI-augmented workplace. ResearchGate. https://www.researchgate.net/publication/392937191_From_Human_to_Hybrid_Leadership_Competencies_in_an_AI_Augmented_Workplace.
- Malthouse, E., & Copulsky, J. (2022). Artificial intelligence ecosystems for marketing communications. *International Journal of Advertising*, 42(1), 128-140.
- Manzoor, A., Qureshi, M. A., Kidney, E., & Longo, L. (2024). A review on machine learning methods for customer churn prediction and recommendations for business practitioners. *IEEE Access*, 12, 70434-70463.
- Marques, P. C. de A., & Oliveira, P. (2024). Integrating artificial intelligence and the balanced scorecard: Strengthening organisational resilience in times of crisis. SSRN. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5050488.
- Marszycki, M. (2020). AI coraz mocniej wkracza w świat e-handlu. ITwiz. <https://itwiz.pl/ai-coraz-mocniej-wkracza-swiat-e-handlu>.
- Martins, J., Cardoso, A., Váz, M., Silva, P., & Abbasi, M. (2025). Performance and Scalability of Data Cleaning and Preprocessing Tools: A Benchmark on Large Real-World Datasets. *Data*, 10(5), 68.
- McGregor, D., & Hostetler, D. (2023). Data-Centric Governance. <https://arxiv.org/abs/-2302.07872>.
- Mengara, M., Avila, J., & Falk, T. (2023). Backdoor attacks to deep neural networks: A survey of the literature, challenges and future research directions. https://www.-researchgate.net/publication/377517071_Backdoor_Attacks_to_Deep_Neural_Networks_A_Survey_of_the_Literature_Challenges_and_Future_Research_Directions.
- Meyer, E., & Bremer, L. (2025). Trust and AI in IT management decision-making: A systematic review and framework for balancing autonomy with human oversight. <https://openreview.net/pdf?id=EEvS0GRiKE>.
- Michałek, J. (2025). Development concept of special economic zones performance. *Polish Journal of Management Studies*, 31, 200-214.
- Mikalef, P., & Gupta, M. (2021). Artificial intelligence capability: Conceptualization, measurement calibration, and empirical study on its impact on organizational creativity and firm performance. *Information & Management*, 58(3), 103434.

- Mumuni, A. (2024). Automated data processing and feature engineering for deep learning and big data applications: a survey. <https://arxiv.org/abs/2403.11395>.
- Musiał, E. (2019). Zastosowania sztucznej inteligencji w obsłudze klientów sklepów internetowych branży mody. Warszawa: Oficyna Wydawnicza SGH.
- Myszak, M., Kokolus, Ł., & Kabaciński, T. (2025). Horyzonty sztucznej inteligencji a przemysł 5.0. Akademia WSB. https://wsb.edu.pl/files/news/13476/sk_horyzonty_sztucznej_inteligencji_a_przemysl_50_ebook_1403.pdf.
- Ndung'u, D. (2022). Data preparation for machine learning modelling. https://www.researchgate.net/publication/361095158_Data_Preparation_For_Machine_Learning_Modelling.
- Nzama, M. L., Epizitone, A., Moyane, S. P., Nkomo, N., & Mthlane, P. P. (2024). The influence of artificial intelligence on the manufacturing industry in South Africa. *South African Journal of Economic and Management Sciences*, 27(1), a5520.
- Okunola, O., Ahsun, M., & Blace, D. (2025). Measuring the ROI of AI-Driven Workforce Transformation Initiatives. *SSRN Electronic Journal*. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5225996.
- Oladele, A. (2025). Cloud-native AI development: Building and deploying scalable machine learning models on AWS, Azure and GCP. https://www.researchgate.net/publication/390704552_Cloud-Native_AI_Development_Building_and_Deploying_Scalable_Machine_Learning_Models_on_AWS_Azure_and_GCP.
- Olak, A. (2025). Badanie EY: Aż 8 na 10 firm z sektora produkcji osiąga swoje cele przy wdrożeniu AI. EY Polska. https://www.ey.com/pl_pl/newsroom/2025/07/ey-ai-przemysl.
- Oluwafemi, I. O., Clement, T., Adanigbo, O. S., Gbenle, T. P., & Adekunle, B. I. (2021). Enhancing customer retention using predictive analytics and personalization in digital marketing campaigns. https://www.researchgate.net/publication/394186053_Enhancing_customer_retention_using_predictive_analytics_and_personalisation_in_digital_marketing_campaigns.
- Oluwafemi, I. O., Clement, T., Adanigbo, O. S., Gbenle, T. P., & Adekunle, B. I. (2021). A review of ethical considerations in AI-driven marketing analytics: Privacy, transparency, and consumer trust (pp. 428–435). *International Journal of Multidisciplinary Research and Growth Evaluation*, 2(2), 428–435.

- Opara-Martins, J., Sahandi, R., & Tian, F. (2016). Critical review of vendor lock-in and its impact on adoption of cloud computing. *Journal of Cloud Computing*, 5(1), 1–18.
- Orr, T. (2025). The state of data & AI engineering 2025. <https://lakefs.io/blog/the-state-of-data-ai-engineering-2025/>.
- Pacheco, P., Rahman, A., Rabbi, F., Fathollahzadeh, M., Abdellatif, T., Shihab, E., Chen, T., Yang, Y., & Zou, Y. (2024). DVC in Open Source ML-development: The Action and the Reaction. https://seal-queensu.github.io/publications/pdf/CAIN_Pouya_2024.pdf.
- Pahune, S., & Akhtar, Z. (2025). Transitioning from MLOps to LLMOps: Navigating the unique challenges of large language models. *Information*, 16(2), 87.
- Paliwal, G., Kumar, A., Sharma, K. P., Bhargava, D., & Shrimal, V. M. (2025). Transformative impact of explainable artificial intelligence: bridging complexity and trust. *Discover Artificial Intelligence*, 5(1), 51.
- Paliwal, M., Kumar, R., Sharma, A., Bhargava, N., & Shrimal, P. (2025). Artificial intelligence for cybersecurity: Trends, challenges, and future directions. *International Journal of Information Security and Privacy*, 15(1), 1–20.
- Pandey, A. C., Gupta, S., & Chhajed, N. (2021). ROI of AI: Effectiveness and measurement. *International Journal of Advanced Research in Computer and Communication Engineering*, 10(6), 749–755.
- Panfil, K. (2025). Zagrożenia sztucznej inteligencji – przykłady i wyjaśnienia: Co musisz wiedzieć w 2025 roku. <https://ttms.com/pl/wyjasnienie-zagrozen-zwiazanych-ze-sztuczna-inteligencja-co-musisz-wiedziec-w-2025-roku/>.
- Parker, P. (2024). What is a customer data platform (CDP) and how does it give marketers the coveted ‘single view’ of their customers? MarTech. <https://martech.org/martech-landscape-customer-data-platform/>.
- Sell, C. (2025). Why customer data platforms need to evolve to meet new industry demands. GrowthLoop Blog. <https://www.growthloop.com/resources/blogs/customer-data-platform-trends>.
- Patel, V., Raut, A., Cheetirala, R., Nadkarni, A., Freeman, R., Glicksberg, B., Klang, E., & Timsina, P. (2024). Cloud Platforms for Developing Generative AI Solutions: A Scoping Review of Tools and Services. <https://arxiv.org/html/2412.06044v1>.

- Patil, D. (2024). Generative artificial intelligence in marketing and advertising: Advancing personalization and optimizing consumer-engagement strategies. SSRN. https://www.researchgate.net/publication/385885224_Generative_artificial_intelligence_in_marketing_and_advertising_Advancing_personalization_and_optimizing_consumer_engagement_strategies.
- Plangger, K., & Campbell, C. L. (2022). Managing in an era of falsity: Falsity from the metaverse to fake news to fake endorsement to synthetic influence to false agendas. *Business Horizons*. <https://doi.org/10.1016/j.bushor.2022.101990>.
- PMR (2025). Polski rynek chmury obliczeniowej pozostaje na fali wzrostowej. <https://www.erp-view.pl/rynek-it/31412-polski-rynek-chmury-obliczeniowej-pozostaje-na-fali-wzrostowej.html>.
- Quinnox (2025). Data Governance for AI: 2025 Challenges, Solutions & Best Practices. <https://www.quinnox.com/blogs/data-governance-for-ai/>.
- Radanliev, P. (2025). AI Ethics: Integrating Transparency, Fairness, and Privacy in AI Development. *Applied Artificial Intelligence*, 39(1), 2463722.
- Rahman, M. (2025). Federated Learning: A Survey on Privacy-Preserving Collaborative Intelligence. <https://arxiv.org/abs/2504.17703>.
- Ramya, J., Swapna, S., Varalakshmi, C., Mukherjee, R., & Kumar, B. R. (2025). AI-Powered Marketing: Predictive Consumer Behavior and Personalized Campaigns. *Journal of Marketing & Social Research*, 2(2), 28–38.
- Ravikumar, S., Sriraman, K., Saketh, T., Lokesh, M., & Karanam, S. (2022). Effect of neural network structure in accelerating performance and accuracy of a convolutional neural network with GPU/TPU for image analytics. <https://pmc.ncbi.nlm.nih.gov/articles/PMC9044238/>.
- Ravikumar, S., Sriraman, K., Saketh, T., Lokesh, M., & Karanam, S. (2022). Effect of neural network structure in accelerating performance and accuracy of a convolutional neural network with GPU/TPU for image analytics. *Cureus*, 14(4), e24069.
- Reda, K., & Jamal, A. (2025). The role of managerial innovation in the achievement of digital transformation: A bibliometric review and critical analysis of the literature. *Alternatives Managériales & Économiques*, 7(3), 378–398.

- Research and Markets (2025). Generative AI cybersecurity market: Generative AI applications, use cases, and growth trends 2025–2030. <https://www.researchandmarkets-.com/reports/5995317/generative-ai-cybersecurity-market-generative-ai>.
- Rikala, P., Braun, G., Järvinen, M., Stahre, J., & Härmäläinen, R. (2024). Understanding and measuring skill gaps in Industry 4.0 — A review. *Technological Forecasting & Social Change*, 201, 123206.
- Robertson, J., Botha, E., Oosthuizen, K., & Montecchi, M. (2025). Managing change when integrating artificial intelligence (AI) into the retail value chain: The AI implementation compass. *Journal of Business Research*, 189, 115198.
- Russell, S., & Norvig, P. (2022). From Animals to Artificial Intelligence: Non-Human Beings' Intellectual Property Protection by "Judicial Capacity for Copyrights". <https://www.scirp.org/journal/paperinformation?paperid=120921>.
- Samola, P. (2025). AI and data lakes: Transforming enterprise data storage and retrieval with machine learning. https://www.researchgate.net/publication/389433288_AI_and_-Data_Lakes_Transforming_Enterprise_Data_Storage_and_Retrieval_with_Machine_Learning.
- Sanderson, C., Douglas, D., & Lu, Q. (2024). Implementing responsible AI: Tensions and trade-offs between ethics aspects. *arXiv*. <https://arxiv.org/pdf/2304.08275>.
- SAP. (2024). Co to jest sztuczna inteligencja dla przedsiębiorstw? <https://www.sap.com/-poland/resources/what-is-enterprise-ai>.
- Sardjono, W., & Pratama, D. A. (2025). Utilizing artificial intelligence predictive maintenance in lean manufacturing to boost industrial sustainability and energy efficiency. *Journal of Theoretical and Applied Information Technology*, 103(13), 4850-4859.
- Sedlak, M., Donta, P., & Pujol, R. (2023). Controlling Data Gravity and Data Friction: From Metrics to Multidimensional Elasticity Strategies. https://www.researchgate.net/-publication/373642785_Controlling_Data_Gravity_and_Data_Friction_From_Metrics_to_Multidimensional_Elasticity_Strategies.

- Sharma, R., Kayalvizhi, K., Maria, H. H., & Suv, B. (2025). Corporate governance and AI ethics: A strategic framework for ethical decision-making in business. *Journal of Information Systems Engineering & Management*, 10(30s), 61–69.
- Shojaei, A., Zameni, M., & Moieni, M. (2025). Anonymity in the Age of AI. <https://www.scirp.org/journal/paperinformation?paperid=144580>.
- Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
- Singh, D., & Gill, S. S. (2023). A systematic review on machine learning and deep learning applications for healthcare. *Journal of Information Security and Applications*, 74, 103513.
- Singh, P. (2023). Systematic review of data-centric approaches in artificial intelligence and machine learning. *Data Science and Management*, 6(3), 144–157.
- Singh, R., & Gill, S. S. (2023). Edge AI: A survey. *Internet of Things and Cyber-Physical Systems*, 3, 71–92.
- Sotonye, F., Shadrach, I., & Chima, K. (2025). Architecting resilient and secure IoT systems: Edge-cloud synergy for offline functionality and regulatory compliance. https://www.researchgate.net/publication/395334032_Architecting_Resilient_and_Secure_IoT_Systems_Edge-Cloud_Synergy_for_Offline_Functionality_and_Regulatory-Compliance.
- Soumya, A. (2025). The Role of Artificial Intelligence and Machine Learning Services in AWS, Google Cloud and Azure. <https://www.ijfmr.com/papers/2025/2/40602.pdf>.
- Stathopoulou, C., Ferikoglou, G., Katsaragakis, M., Masouros, C., Xydis, S., & Soudris, D. (2025). SynergAI: Edge-to-Cloud Synergy for Architecture-Driven High-Performance Orchestration for AI Inference. <https://arxiv.org/abs/2509.12252>.
- Stelmach, A. (2025). Przegląd współczesnych tendencji doboru i wdrażania systemów informatycznych oraz rozwiązań opartych na sztucznej inteligencji w polskim przemyśle. <http://bazekon.icm.edu.pl/bazekon/element/bwmeta1.element.ekon-element-000171714755>.

- Stone, E., Patel, A., Ghiasi, M., Mittal, P., & Rahimi, A. (2025). Navigating MLOps: Insights into Maturity, Lifecycle, Tools, and Careers. <https://arxiv.org/html/2503.-15577v1>.
- Sun, Y., Liu, C., & Gao, Y. (2023). Artificial intelligence for scientific discovery and design: Past, present, and future. *Frontiers in Pharmacology*, 14, 9958434.
- Ślusarczyk, B. (2019). Potencjalne rezultaty wprowadzania koncepcji przemysłu 4.0 w przedsiębiorstwach. *Przegląd Organizacji*, 1(948), 4-10.
- Ślusarczyk, Z. (2021). Znaczenie sztucznej inteligencji w działalności przedsiębiorstw. Podstawowe informacje. *Zarządzanie Innowacyjne w Gospodarce i Biznesie*, 2(33), 49–58.
- Tarnas, P. (2025). Lider w erze AI: jak zachować ludzką przewagę w świecie algorytmów. <https://mitsmr.pl/przywodztwo/przywodztwo-w-erze-ai-jak-budowac-przewage/>.
- Techreviewer. (2025). AI in software development 2025: From exploration to accountability – a global survey-based analysis. <https://techreviewer.co/blog/ai-in-software-development-2025-from-exploration-to-accountability-a-global-survey-analysis>.
- Tewari, S. (2025). AI Powered Data Governance - Ensuring Data Quality and Compliance in the Era of Big Data. *New Journal of Artificial Intelligence and General Science*, 8(1), 187–197.
- Themann, L. (2025). Challenges and Strategies Used In Implementing AIGovernance: A Systematic Literature Review. Stockholm University. <https://su.diva-portal.org/smash/-get/diva2:1983756/FULLTEXT01.pdf>.
- Thota, R. (2024). Optimizing edge computing and AI for low-latency cloud workloads. https://www.researchgate.net/publication/389905403_Optimizing_edge_computing_and_AI_for_low-latency_cloud_workloads.
- Tjondronegoro, D. (2024). Strategic AI Governance: Insights from Leading Nations. arXiv. <https://arxiv.org/abs/2410.01819>.
- Toorajipour, R., Oghazi, P., & Palmiś, M. (2024). Data ecosystem business models: Value propositions and value capture with Artificial Intelligence of Things. *International Journal of Information Management*, 78, 102804.

- Triguero, I., Molina, D., Poyatos, J., Del Ser, J., & Herrera, F. (2023). General-Purpose Artificial Intelligence Systems (GPAIS): Properties, Definition, Taxonomy, Societal Implications and Responsible Governance. arXiv. <https://doi.org/10.48550/arXiv.-2307.14283>.
- Tsitsoula, S. (2025). The role of artificial intelligence in shaping strategic business decision-making (MSc thesis, Hellenic Open University). <https://apothesis.eap.gr/-archive/download/695cd8e2-ccd0-4029-9955-ecf53abf7b35.pdf>.
- Tursunbayeva, A., & Gal, U. (2024). Artificial intelligence in organizations: A review and research agenda from a socio-technical perspective. *Business Horizons*, 67(4), 101716.
- Uchwat, M. (2025). State of the Polish AI 2025: Rozwój sztucznej inteligencji w Polsce – ogromny potencjał i wyzwania. <https://tjsoft.pl/state-of-the-polish-ai-2025-rozwoj-sztucznej-inteligencji-w-polsce-ogromny-potencjal-i-wyzwania/>.
- Uren, V., & Edwards, J. S. (2023). Technology readiness and the organizational journey towards AI adoption: An empirical study. *International Journal of Information Management*, 68, 102588.
- Utama, H. K., & Sari, R. (2025). The influence of strategy implementation on competitive advantage. *Jurnal Manajemen Perbankan Keuangan Nitro*, 1(3), 25–36.
- Vargas, J., Maldonado, A., & Vairetti, C. (2023). A predict-and-optimize approach to profit-driven churn prevention. arXiv. <https://arxiv.org/html/2310.07047v2>.
- Vijona, L. (2025). An introduction to GPU performance optimization for deep learning: Key techniques & tools. <https://www.centron.de/en/tutorial/an-introduction-to-gpu-performance-optimization-for-deep-learning-key-techniques-tools/>.
- Vogel, A. T., Vogel, J., Watchravesringkan, K., Cook, S. C., Beasley, J., Croom, R., Peterson, D., & Finkelstein, J. (2023). Design piracy: An interdisciplinary investigation into competitive industrial behavior. *Journal of Business Research*, 164, 113946.
- Wang, J., & Luo, J. (2022). A Review of the Optimal Design of Neural Networks Based on FPGA. *Applied Sciences*, 12(21), 10771.

- Weinzierl, S., Dunzer, S., Zilker, S., & Matzner, M. (2020). Prescriptive business process monitoring for recommending next best actions. In E. Fahland, C. Ghidini, D. Becker, & M. Dumas (Eds.), *Business Process Management Forum (BPM 2020)* (pp. 193–209). Springer. https://doi.org/10.1007/978-3-030-58638-6_12.
- William, A. (2025). Machine learning algorithms for demand forecasting. ResearchGate. https://www.researchgate.net/publication/389357099_Machine_Learning_Algorithms_for_Demand_Forecasting.
- Wu, M., Liu, C., & Huang, J. (2021). The special economic zones and innovation: Evidence from China. *China Economic Quarterly International*, 1(4), 319–330.
- Xie, Y., Chen, Y., Wei, Q., & Yin, H. (2024). A hybrid deep learning approach to improve real-time effluent quality prediction in wastewater treatment plant. *Water Research*, 250, 121092.
- Xu, G., Wang, Z., Li, J., Wang, Y., Zhao, Q., Chen, H., Yu, X., Liu, D., & Wang, J. (2025). Large Language Models for Cyber Security: A Systematic Literature Review. <https://arxiv.org/html/2405.04760v5>.
- Yao, Y., Zhang, T., Yao, X., Wang, X., & Ma, J. (2021). Training deep neural network in limited data scenarios by transfer learning. <https://arxiv.org/abs/2111.06061>.
- Yu, D., Yang, B., Liu, D., Wang, H., & Pan, S. (2023). A survey on neural-symbolic learning systems. arXiv:2111.08164. <https://arxiv.org/abs/2111.08164>.
- Zeng, D. Z. (2021). The past, present, and future of special economic zones and their impact. *Journal of International Economic Law*, 24(2), 259–275.
- Zhao, P., Zhu, W., Jiao, P., Gao, D., & Wu, O. (2025). Data poisoning in deep learning: A survey. <https://arxiv.org/pdf/2503.22759>.
- Zreikat, A., AlArnaout, M., Abadleh, A., Elbasi, L., & Mostafa, S. (2025). The Integration of the Internet of Things (IoT) Applications into 5G Networks: A Review and Analysis. *Computers*, 14(7), 250.

Monografia „Gotowość technologiczna i wykorzystanie sztucznej inteligencji w przedsiębiorstwach: analiza na przykładzie Specjalnych Stref Ekonomicznych” stanowi pogłębione studium uwarunkowań organizacyjnych, technologicznych i strategicznych determinujących proces wdrażania sztucznej inteligencji w przedsiębiorstwach funkcjonujących w polskich Specjalnych Strefach Ekonomicznych. Autor podejmuje kompleksową analizę infrastruktury danych, architektur systemowych, kompetencji zarządczych oraz zastosowań AI w kluczowych obszarach funkcjonalnych organizacji.

Praca łączy szeroką panoramę teoretyczną z empiryczną oceną dojrzałości cyfrowej 177 przedsiębiorstw, weryfikując zależności między poziomem adopcji rozwiązań AI a osiąganymi wynikami biznesowymi. Szczególną wartość stanowi identyfikacja czynników moderujących i mediujących wpływ AI na efektywność organizacyjną, a także analiza roli kierownictwa i kultury innowacji w procesie transformacji cyfrowej.

Monografia wnosi istotny wkład w dyskusję nad gotowością technologiczną przedsiębiorstw, przedstawiając spójne ramy koncepcyjne oraz wyniki badań o wysokiej użyteczności poznawczej i aplikacyjnej. Publikacja adresowana jest do badaczy, praktyków zarządzania oraz decydentów publicznych zainteresowanych rozwojem kompetencji cyfrowych, polityką innowacyjną oraz implementacją sztucznej inteligencji w warunkach polskiej gospodarki.