

Piotr Morciniec
Instytut Nauk Teologicznych, Uniwersytet Opolski
<https://orcid.org/0000-0001-6312-8296>

Etyka sztucznej inteligencji – fakty i perspektywy

Ethics of Artificial Intelligence – facts and perspectives

Abstract

Research context: The dynamic development of artificial intelligence (AI) has revolutionised the modern era, bringing with it both enormous benefits and increasing ethical and social risks. A moral reflection on its applications, implications, and consequences is therefore justified.

Research objective: The main aim of the article is to establish moral frameworks for the functioning of AI by distinguishing human intelligence from machine intelligence and by analysing who actually bears ethical responsibility for AI's actions, thereby identifying the subject of moral responsibility.

Research method: The study employs desk research, i.e., an analysis of existing data, including selected documents and studies addressing ethical aspects of AI use, with particular attention given to the activity in this field of the pope of the “digital revolution,” Leo XIV. Comparative analysis was also used: ethical principles of AI identified by a selected chatbot were juxtaposed with definitions of morality and ethics, and approaches to the issue were compared in the Vatican's *Nota Antiqua et nova* and the EU's AI Act.

Results: It was established that AI cannot act morally because it is not a human subject (a person), which means that programmers and users, not machines, ultimately bear responsibility for its use and any potential harm.

Conclusions: It is necessary for the development of AI to place human dignity at its center and to be grounded in strong ethical frameworks and legal regulations, so that the benefits of this technology outweigh the risks, such as deepening loneliness, undermining the right to work, autonomous AI decisions in military actions, or the educational danger of questioning the rationality of the human subject.

Keywords: ethics of artificial intelligence; dignity of the human person; moral responsibility; anthropomorphization of machines; Nota Antiqua et nova; AI Act; digital revolution; threats associated with AI.

Abstrakt

Kontekst badań: Dynamiczny rozwój sztucznej inteligencji (AI) zrewolucjonizował epokę, niosąc ze sobą zarówno ogromne korzyści, jak i narastające zagrożenia etyczne i społeczne. Zasadny jest więc namysł moralny nad jej zastosowaniami, implikacjami i konsekwencjami.

Cel badań: Głównym celem artykułu jest ustalenie ram moralnych dla funkcjonowania AI, poprzez odróżnienie inteligencji ludzkiej od maszynowej oraz analizę, kto faktycznie ponosi odpowiedzialność etyczną za jej działania, a więc zidentyfikowanie podmiotu odpowiedzialności moralnej.

Metoda badawcza: W badaniach zastosowano desk research, a więc analizę danych zastanych, z uwzględnieniem wybranych dokumentów i opracowań dotyczących etycznych aspektów stosowania AI, szczególną uwagę poświęcając aktywności na tym polu papieża „rewolucji cyfrowej”, Leona XIV. Posłużono się także analizą porównawczą, zestawiając zasady etyczne AI wskazane przez wybranego chatbota z definicjami moralności i etyki oraz porównując podejścia do problemu w watykańskiej Nocie Antiqua et nova i unijnym AI Act.

Osiągnięte wyniki: Ustalono, że AI nie może działać moralnie, ponieważ nie jest ludzkim podmiotem (osobą), co oznacza, że to programiści i użytkownicy, a nie maszyny, ponoszą ostateczną odpowiedzialność za jej użycie i potencjalne szkody.

Wnioski: Konieczne jest, aby rozwój AI stawiał w centrum godność człowieka i opierał się na silnych ramach etycznych oraz regulacjach prawnych, aby dobro wynikające z tej technologii przewyższało zagrożenia, takie jak pogłębianie samotności, zakwestionowanie prawa do pracy, autonomiczne decyzje AI w działaniach zbrojnych czy niebezpieczeństwo edukacyjnego zakwestionowania racjonalności ludzkiego podmiotu.

Słowa kluczowe: etyka sztucznej inteligencji, godność osoby ludzkiej, odpowiedzialność moralna. antropomorfizacja maszyn, Nota Antiqua et nova, AI Act, rewolucja cyfrowa, zagrożenia związane z AI.

Sztuczna inteligencja zrewolucjonizowała naszą epokę. Postęp technologiczny sprawia, że AI staje się coraz bardziej dostępna, a wręcz powszechna. Sztuczna inteligencja wychodzi naprzeciw potrzebom przeciętnego człowieka w różnych obszarach codzienności, zapewniając coraz większą efektywność, wygodę i personalizację oraz pomagając zachować zdrowie i bezpieczeństwo (Ilnicka, 2024, p. 29). Tak pozytywne ujęcie powinno jednak uwzględnić i dru-

gą stronę medalu, tzn. fakt, że wraz z eksplozją innowacyjności narastać będą także zagrożenia i wyzwania. Nie ulega bowiem wątpliwości, że zasadne jest mówienie o ambiwalentnym charakterze dynamicznego rozwoju zastosowań sztucznej inteligencji (Bugajski, 2023). Nic więc dziwnego, że w tomie dedykowanym sztucznej inteligencji nie sposób pominąć kwestii etycznych, a więc pytań o dobro i zło tych przełomowych osiągnięć ludzkiej myśli (Sypniewska and Gołębiowski, 2023).

Punktem wyjścia należy jednak uczynić opis stanu faktycznego, a więc świadomość specyficznej bezradności wobec problemów i zagadnień całkowicie nowych oraz w dużej mierze odmiennych od całej dotychczas podejmowanej problematyki moralnej. Już konfrontacja z rzeczywistością wirtualną i przestrzenią alternatywną pokazuje, jak nowe modele i sposoby podejścia konieczne są w etycznym namyśle nad tymi nowymi zjawiskami. Na początek warto więc zebrać podstawowe fakty dotyczące sztucznej inteligencji. Sztuczna inteligencja niewątpliwie wpływa na osobę ludzką, kształtuje ją i zmienia (Kuciński, 2021). Ten fakt należy zestawić z drugim, z tym mianowicie, że kierunek tych zmian jest niewiadomy, nie do końca możliwy do przewidzenia, natomiast ich dynamika wciąż zadziwia i w jakieś mierze niepokoi.

Nawet niewielka wiedza na temat AI pozwala stwierdzić, że chatbot zawsze udziela odpowiedzi i nie zna odpowiedzi „nie wiem”. Już na podstawie tej prostej konstatacji można wyprowadzić wniosek etyczny, że w przypadku, kiedy sztuczna inteligencja nie zna odpowiedzi, będzie ją konfabulować (halucynacje AI); (Szpyrka, 2025). W kategoriach moralnych mówilibyśmy wtedy o kłamstwie albo o komunikowaniu informacji nieprawdziwych, ale w przypadku tych skomplikowanych maszyn jest to język nieadekwatny. Niewątpliwie jednak mamy do czynienia z dezinformacją tworzoną przez generatywne systemy AI (Wyrostkiewicz, 2025, pp. 369–374). Użycie słowa „maszyna” każe w ramach analizy sztucznej inteligencji użyć kategorii „robot”, która pozwala sformułować proste a znaczące równania: człowiek nie równa się maszyna, maszyna nie równa się człowiek. Owa niemożność postawienia znaku równości między osobą ludzką a sztuczną inteligencją gra ważną rolę w formułowaniu zasad etyki AI.

Po tym merytorycznym wstępie przejdźmy do zasadniczego tematu niniejszego artykułu. Pytając o ramy moralne funkcjonowania sztucznej inteligencji, trudno nie ulec pokusie, żeby zapytać konkretnego chatbota o wiedzę, czyli informacje, jakie zasady moralne obowiązują w działaniu sztucznej inteligencji. To pytanie autor postawił czatbotowi Gemini. Uzyskane informacje okazały się precyzyjne, poprawne, ale od razu pojawiło się pytanie, czy same tak sformułowane normy są wystarczające dla etycznie poprawnego funkcjonowania sztucznej inteligencji. Zestaw wskazanych reguł był następujący:

1. *Przejrzystość i wyjaśnialność (Transparency and Explainability)*: Należy dążyć do tego, aby procesy decyzyjne sztucznej inteligencji były zrozumiałe dla człowieka. Użytkownik powinien wiedzieć, dlaczego AI podjęła taką, a nie inną decyzję lub wydała taki właśnie werdykt. Nie może być to „czarna skrzynka” (*black box*).
2. *Odpowiedzialność (Accountability)*: Musi istnieć jasna odpowiedzialność za działania systemów AI. Jeśli system wyrządzi szkodę, powinna istnieć możliwość wskazania osoby lub organizacji, która jest za to odpowiedzialna.
3. *Bezpieczeństwo i niezawodność (Safety and Reliability)*: Systemy AI muszą być projektowane w taki sposób, aby były bezpieczne, niezawodne i nie powodowały nieprzewidzianych szkód. Trzeba minimalizować ryzyko błędów i awarii.
4. *Sprawiedliwość i brak uprzedzeń (Fairness and Non-bias)*: Algorytmy AI nie mogą dyskryminować ludzi ze względu na płeć, rasę, pochodzenie etniczne, wiek czy inne cechy. Należy eliminować uprzedzenia zawarte w danych treningowych, aby AI działała sprawiedliwie dla wszystkich.
5. *Prywatność (Privacy)*: Systemy AI często przetwarzają ogromne ilości danych osobowych. Muszą one przestrzegać zasad ochrony prywatności, a dane użytkowników powinny być bezpiecznie przechowywane i wykorzystywane wyłącznie w celach, na które użytkownik wyraził zgodę.
6. *Ludzka kontrola (Human Control)*: Nawet w zaawansowanych systemach AI, człowiek powinien mieć możliwość nadzoru i interwencji. AI ma być narzędziem, które wspiera ludzi, a nie całkowicie ich zastępuje w podejmowaniu kluczowych decyzji.
7. *Dobrostan społeczny i środowiskowy (Social and Environmental Well-being)*: Rozwój i wdrażanie AI powinny przyczyniać się do poprawy jakości życia społeczeństwa i nie szkodzić środowisku naturalnemu. Należy myśleć o długoterminowych konsekwencjach technologii.

Usystematyzowane informacje przekazane przez sztuczną inteligencję pozwalają metodą przybliżeń podjąć temat zasadniczy. Najpierw jednak należy odnotować, że istnieją też ślepe zaułki intelektualne, w które nie warto się zapuszczać w poszukiwaniu rozwiązań. Jeden z nich ukazał prof. Charles Camosy, bioetyk z katolickiego Uniwersytetu Ameryki, stwierdzając w wywiadzie telewizyjnym w odniesieniu do AI: „Nie jest jeszcze za późno, by zamknąć [tego] dzina z powrotem w butelce i uniknąć najgorszych skutków nowego wynalazku” (eKAI.pl, 2025). Samo wyobrażenie o tym, że można odwrócić scenariusz rozwoju sztucznej inteligencji, aby zapobiec domniemanym szkodom spowodowanym

przez ten wynalazek, należy uznać za niedorzeczne. Zamiast roztaczać katastroficzne wizje, trzeba skupić się na takim kształtowaniu traktu rozwojowego AI, aby dobro, które można wyprowadzić z tej potężnej siły, przeważało nad potencjalnymi zagrożeniami.

Skoro nie zamierzamy demonizować sztucznej inteligencji, zapytajmy o samą możliwość „nauczenia jej” moralności. Z pomocą zdaje się przychodzić wypowiedź jednego z popularyzatorów nauki, dra Tomasza Rożka, który w roku w 2023 stwierdził: „Próbowano nauczyć sztuczną inteligencję (AI) moralności i to zawsze wychodziło źle” (Rasińska, 2023). Autor tej wypowiedzi wyjaśniał dalej: „Bez wątpienia muszą powstać regulacje dotyczące sztucznej inteligencji, ale nie sądzę, aby ktokolwiek miał pomysł, jak powinny one wyglądać. Będąc na początku jakiejś drogi, wiele rzeczy możemy sobie wyobrazić, natomiast wiedza to zupełnie inna kwestia”. Współczesność pokazała, iż takie regulacje powstają, ale problem nie leży w samych normach czy regulacjach. Błądność założenia, że AI można nauczyć moralności, ujawnia się, kiedy zapytamy o to, czym jest moralność i porządkująca ją nauka – etyka.

„Moralność to zespół zasad, norm i wartości, które określają, co jest dobre, a co złe w zachowaniu i postawach jednostki (ludzkiej) lub grupy (osób)” – odpowiedział na pierwsze pytanie chatbot. A sięgnięcie do podręcznika etyki uzupełnia tę wiedzę o definicję, według której: „Etyka to nauka filozoficzna, która ustala moralne podstawy i reguły ludzkiego działania przy pomocy wrodzonych człowiekowi zdolności poznawczych” (Ślipko, 1974, p. 15).

W obydwu określeniach mowa jest o działaniu ludzkim, które charakteryzuje się świadomością (rozumnością) i wolnością. Tylko takie działanie może być przypisane podmiotowi i za takie działanie podmiot odpowiada. Problem w tym, że takim rozumnym i wolnym działającym podmiotem, który ponosi odpowiedzialność za działanie, jest człowiek, a nie najbardziej nawet rozwinięta i skomplikowana maszyna. Sztuczna inteligencja zdecydowanie wymienionych wymagań nie spełnia, nie może zatem działać moralnie, a próby „nauczenia jej moralności” byłyby same w sobie sprzeczne z logiką i skazane na porażkę. Powyższe konstatacje nakazują więc sformułować podstawowe i kluczowe dla dalszego rozumowania rozróżnienia.

1. Fundamentalne rozróżnienia

Faktem jest, że wiele zastosowań AI budzi wątpliwości natury etycznej, które zyskują na znaczeniu w praktycznych aplikacjach, jednak to nie chatboty są adresatami wymagań moralnych. Jak bowiem „wprowadzić wymiar etyczny do

algorytmu (czyli wzoru matematycznego), który nie ma serca, duszy, umysłu?” – pyta jeden ze znawców problematyki (Lennox, 2023, p. 21). Autor proponuje najpierw odróżnić wąską (słabą) AI, zaprogramowaną do wykonywania jednego zadania, od generatywnej (mocnej) AI, która ma niewyobrażalnie większe możliwości techniczne. O wiele jednak istotniejsze od tego rozróżnienia jest wystrzeżenie się antropomorfizacji maszyn, którym w potocznej, a jeszcze bardziej w medialnej narracji przypisuje się pojęcia właściwe osobie ludzkiej, takie jak planowanie, odczuwanie, rozumowanie, przeżywanie. Jako imperatyw moralny należy potraktować rewizję narracji, która pojęcia służące opisowi ludzkich zachowań przypisuje nieożywionym maszynom. Taki zabieg jest groźny społecznie w swoim przesłaniu, a w przekazie medialnym potęguje sensacyjność doniesień, obciążając je nadmiernym dramatyzmem, przesadnym optymizmem lub obawami budzącymi bezrefleksyjny strach. Media żywią się strachem przed sztuczną inteligencją, każdą pomyłkę przypisując AI, a nie ludziom, którzy często zadają mało precyzyjne pytania i nie weryfikują źródeł informacji.

Ciekawą myśl na temat antropomorfizacji robotów wyraził jeden z naukowców, ostrzegając: „Należy zachować ostrożność, przyjmując założenie, że ludzkość ma intelektualną zdolność stworzenia inteligencji rywalizującej z inteligencją ludzką. [...], niezależnie od tego, jaką ilością czasu będziemy dysponować” (Lennox, 2023, p. 21). Z kolei pewna badaczka nazwała po imieniu utopijne i błędne przekonanie, które określiła jako „wiarę”, „że możemy stworzyć, co chcemy” (Spiekermann *et al.*, 2025). Na takich założeniach i przekonaniach bazują mechanizmy demonizujące AI i sterujące generalizacją lęków.

Ostrzeżenie przed narracją antropomorficzną zasadza się na faktach. Sztuczna inteligencja jest zbiorem coraz doskonalszych algorytmów, zmieniających się co miesiąc. Rozwiązania użyte w modelach otwartych, dających pełny dostęp do ich kodu, lub w modelach częściowo otwartych, które udostępniają parametry swoich modeli, konkurują z wielkimi bardzo kosztownymi modelami, których szczegóły nie są dostępne. Jest wiele nowych pomysłów, które mogą zastąpić dominującą obecnie architekturę transformerów (GPT = Generative Pre-trained Transformer oznacza: generatywne wstępnie wytrenowane transformery). [...] Modele OpenAI (GPT-5, o3, GPT-4o), Claude firmy Anthropic, Gemini Googla, oraz Grok-4 xAI to modele zamknięte, do których parametrów ani kodu nie mamy dostępu (Duch, 2025). Problem systemów zamkniętych jest sam w sobie problemem etycznym, ponieważ między impulsem wysłanym do chatbota (*input*) a uzyskanym produktem matematycznym (*output*) pojawia się „czarna skrzynka” (*black box*), czyli maszyna z nieznanymi klientowi algorytmami, które w różny sposób modyfikują (kształtują) odpowiedź. W tym kontekście werbalizowany jest temat wewnętrznej cenzury i „biasów”, czyli tendencyjności, stronniczości i uprze-

dzeń generowanych przez odpowiednio zapisane algorytmy i w efekcie tendencyjnie gromadzone dane (Bierzyński, 2025a). Powyższe treści odsyłają do wątku „ludzkiego” pochodzenia problemów i zjawisk nieetycznych, bez obciążania winą maszyn liczących.

Siła nowoczesnych systemów AI wynika z dużej mocy obliczeniowej, jaką dysponują one w celu opanowania ogromu danych, profilowania jednostek i wykrywania wzorców (schematów). Między inteligencją maszynową a ludzką istnieją jednak fundamentalne różnice, których nie przewyżczy żadna liczba badań. Bardzo trafny jest w tym kontekście tytuł wystąpienia prof. J. Mc Mellichamp w ramach konferencji „Artificial Intelligence and Human Mind (Yale University, już w 1986): „Wyraz «sztuczna» w określeniu «sztuczna inteligencja» ma jak najbardziej realne znaczenie”. We współczesnych debatach termin ten wydaje się kluczowy. Znakiem ludzkiej inteligencji jest zdolność wizualizowania rzeczy i myślenia o scenariuszach zawierających przedmioty i procesy istniejące jedynie w umyśle. Ta zdolność jest fundamentalna u człowieka i odróżnia nas od maszyn (Crookes, za: Lennox, pp. 23–24). Potwierdza to również ludzka zdolność do konstruktywnego myślenia – w tym kontekście: maszyna nie stworzy czegoś nowego, w przeciwieństwie do potęgi ludzkiej inteligencji. Ten wątek domaga się oddzielnego, dogłębnego opracowania, co warto będzie podjąć, łącząc treści teoretyczne z analizą przesłania tematycznego dokumentu watykańskiego, *Antiqua et nova*.

Dotychczasowy tok rozumowania pozwala na wyciągnięcie pierwszego wniosku etycznego. Ponieważ komputerowe systemy AI nie są osobami, nie posiadają sumienia, więc moralny charakter ich „decyzji” będzie odzwierciedlać moralność ich programistów oraz użytkowników (Lennox, 2023, p. 116; Bierzyński, 2025b), a nie „moralność postępowania maszyny”. Wniosek ten falsyfikuje samą już sugestię, jakoby istniała „etyka sztucznej inteligencji”. A wątpliwości, jakie się jednak wokół tej kwestii pojawiają, dotyczą moralności działania wspomnianych wyżej ludzkich podmiotów, nie zaś maszyn. Poprawnie więc należy mówić o etyce programowania i korzystania z obliczeniowych możliwości maszyn, a tytuł artykułu powinien brzmieć: „Etyka wobec sztucznej inteligencji”. Warto jednak dodać, że im większa będzie „swoboda działania” maszyn, zwłaszcza przy minimalizacji ludzkiej kontroli tego specyficznego rachunku prawdopodobieństwa, tym więcej standardów moralnych będzie potrzebnych do ich bezpiecznego działania. Wiele obszarów wykorzystania AI w praktyce potwierdza taką tezę i takie niebezpieczeństwa, np.:

- Selekcja kandydatów na rynku pracy i dyskryminacja
- Pojazdy autonomiczne a odpowiedzialność
- Systemy rozpoznawania twarzy a prywatność i inwigilacja

- Broń autonomiczna a nowoczesna wojna
- Nachalna reklama oparta na śledzeniu nawyków konsumenta a wolność wyboru (decyzji) itd.

Bardzo trafnie spointował wątek pytania o aspekty etyczne związane ze sztuczną inteligencją prof. Manfred Spitzer, prowadzący od 30 lat badania nad sieciami neuronowymi, które leżą u podstaw AI. Stwierdził on: „Nie musimy myśleć o tym, czy ogólna sztuczna inteligencja w końcu zyska świadomość albo czy zniszczy ludzkość – wszystko to science fiction. Musimy jednak dogłębnie zastanowić się nad realnym ryzykiem i niebezpieczeństwami związanymi z jej nadużywaniem przez ludzi o złych zamiarach do realizacji własnych celów. A także nad odpowiedzialnością najbogatszych przedsiębiorstw na świecie” (Spitzer, 2025, p. 295).

2. Leon XIV wobec sztucznej inteligencji

Może wydawać się zaskakujące, że po przedstawieniu fundamentalnych założeń dotyczących etyki sztucznej inteligencji od razu przywołuję postać wybranego w tym roku papieża, Leona XIV. Nie jest to jednak zabieg nieuzasadniony. Kiedy Papież przyjął imię Leona XIV – nawiązując do roli swojego poprzednika, Leona XIII, autora kluczowej na swoje czasy encykliki „*Rerum novarum*”, stało się jasne, że świadomy jest odpowiedzialności, jaka ciąży na przywódcach Kościoła katolickiego. Leon XIII pod koniec XIX wieku przekazał światu katolicką odpowiedź na przemiany społeczne i ekonomiczne epoki rewolucji przemysłowej. W dobie sztucznej inteligencji i rewolucji cyfrowej Kościół staje wobec nowej – tym razem cyfrowej – rewolucji przemysłowej (określanej jako „czwarta rewolucja przemysłowa”, zob. Skalfist, Mikelsten and Teigens, 2020, p. 375) i ma udzielić odpowiedzi na pytanie: jak chronić godność człowieka w świecie radykalnie nowych technologii? Sam nowo wybrany papież stwierdził: „Tak jak industrializacja przyniosła dobrobyt, ale i wyzysk, tak dziś rozwój algorytmów i systemów językowych (jak ChatGPT) wymaga od Kościoła stanowczego głosu w obronie człowieka, jego praw, pracy i tożsamości”. Rozwój sztucznej inteligencji jest bowiem nie tylko wyzwaniem technologicznym, ale też powołaniem duchowym. „Nie chodzi tylko o to, co AI może zrobić, lecz kim stajemy się przez technologie, które budujemy” – napisał Leon XIV w przesłaniu do uczestników Builders AI Forum 2025 w Rzymie (Leo XIV, 2025).

Kluczową rolę Leona XIV doceniono krótko po jego wyborze. „Time Magazine”, który w 2023 roku zainicjował coroczną listę osób najbardziej znaczących

dla AI, umieścił papieża Leona XIV na liście „Time 100 AI” na rok 2025 w kategorii: „Myśliciel” (Shah, 2025).

Celem listy jest „pokazanie, że kierunek rozwoju sztucznej inteligencji nie jest wyznaczany przez maszyny, ale przez ludzi – innowatorów, zwolenników, artystów i wszystkich zaangażowanych w przyszłość tej technologii” (Jacobs, 2025). Uzasadniając nominację papieża, pisano o jego dokonaniach ekologicznych, ale również zacytowano jego apel do obywateli, organizacji pozarządowych i grup nacisku, aby wywierać presję na rządy, w celu ograniczeniu szkód wyrządzonych środowisku. Natomiast katoliccy komentatorzy dostrzegli także potencjał regulacyjny w odniesieniu do sztucznej inteligencji: „Jego wybór imienia jest hołdem dla Leona XIII, który sprawował urząd podczas rewolucji przemysłowej pod koniec XIX wieku i sprzeciwiał się nowym, napędzanym przez maszyny systemom gospodarczym, które zamieniały pracowników w towary [...] jeśli Leon XIV będzie nadal uwrażliwiał 1,4 mld katolików na «alienujący potencjał» sztucznej inteligencji, pojawi się «ważna i nieoczekiwana duchowa przeciwwaga dla Doliny Krzemowej»” (eKAI, 2025).

Z perspektywy katolickiej Kościół pragnie w obliczu dynamicznego rozwoju sztucznej inteligencji wnieść spokojny i głęboko refleksyjny głos do debaty nad jej etycznym wykorzystaniem. Potwierdził to Leon XIV w przesłaniu na II Rzymską Konferencję o AI (06.2025), (Darmoros, 2025), deklarując, że Kościół chce „uczestniczyć w tych rozmowach, które bezpośrednio dotyczą teraźniejszości i przyszłości naszej ludzkiej rodziny. [...] Sztuczna inteligencja, zwłaszcza generatywna AI, otworzyła nowe horyzonty na wielu płaszczyznach – w tym w zakresie rozwoju badań w ochronie zdrowia i odkryć naukowych – ale rodzi również niepokojące pytania związane z jej możliwym wpływem na otwartość ludzkości na prawdę i piękno oraz na naszą wyjątkową zdolność pojmowania i przetwarzania rzeczywistości”. Papież ostrzegał przed myleniem ilości informacji z prawdziwą inteligencją: „Dostęp do danych – nawet najbardziej rozległy – nie może być utożsamiany z inteligencją, która zakłada otwartość człowieka na ostateczne pytania życia i ukierunkowanie na Prawdę i Dobro”. Identyfikując obszary newralgiczne, Leon XIV dał wyraz szczególnemu z troskaniu o młode pokolenie, wskazując na możliwe konsekwencje korzystania z AI dla ich rozwoju neurologicznego i moralnego. „Są oni naszą nadzieją na przyszłość, a dobro społeczeństwa zależy od tego, czy zostanie im dana możliwość rozwijania ich darów i zdolności otrzymanych od Boga oraz odpowiadania na wyzwania współczesności i potrzeby innych z wolnym i ofiarnym duchem”. Podejmując wątek integralnego rozwoju człowieka i społeczeństw, papież wyjaśnił, że chodzi o dobro osoby ludzkiej nie tylko w wymiarze materialnym, ale także intelektualnym i duchowym. AI może być używana do wspierania równości, ale równie łatwo może służyć in-

teresom nielicznych, „a nawet – co gorsza – podsycać konflikty i przemoc”. Do tego samego wymiaru nawiązał w słowie na Builders AI Forum 2025, w którym życzył uczestnikom: „sztuczna inteligencja, jak każdy ludzki wynalazek, wyraża z kreatywności, którą Bóg powierzył człowiekowi. [...] Niech wasza współpraca zaowocuje AI, która odzwierciedla Boży zamysł – inteligentną, relacyjną i kierowaną miłością” (Wiara.pl, 2025).

Pozostając w zgodzie z faktami, należy dodać, że już papież Franciszek akcentował etyczne wymagania związane z korzystaniem ze sztucznej inteligencji, a rok 2025 okazał się przełomowy, gdy idzie o apelatywne dokumenty o dużej wadze merytorycznej, ogłoszone zarówno powagą Kościoła katolickiego, jak i Unii Europejskiej. Chodzi o dwa dokumenty (Omówienie porównawcze: Salar, 2025):

Antiqua et nova. Nota na temat relacji pomiędzy sztuczną inteligencją a inteligencją ludzką, opublikowany 28 stycznia, jako owoc współpracy między Dykasterią Nauki Wiary a Dykasterią ds. Kultury i Edukacji (AeN) (Dykasteria Nauki Wiary i Dykasteria ds. Kultury i Edukacji, 2025) oraz

Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 z dnia 13 czerwca 2024 r. ustanawiające zharmonizowane zasady w zakresie sztucznej inteligencji, którego zakazy i przepisy ogólne weszły w życie 2 lutego 2025 (AI Act) (Eur-lex.europa.eu, 2024).

Należy odnotować ten fakt ze świadomością, że same regulacje prawne czy polityczne są niewystarczające. Konieczna jest pozytywna dynamika instytucji, społeczności i jednostek, które kształtują właściwy klimat etyczny. Warto jednak zestawzić ujęcie katolickie z cywilnoprawnym (unijnym), aby wyartykułować podobieństwa i różnice.

3. Istotne akcenty w *Antiqua et nova* i w AI Act

Pełne i wyczerpujące omówienie merytorycznej zawartości obydwu dokumentów wymagałoby obszernego opracowania. Mając świadomość skrótów myślowych i uproszczeń, spróbujemy jednak przynajmniej hasłowo zebrać przesłanie tych enuncjacji.

Obydwa dokumenty podejmują kwestie związane z AI z różnych perspektyw. Zauważa się, że wspólne są kluczowe zdiagnozowane problemy i uzupełniające się ramy. *Antiqua et nova* ma na celu pozytywne rozeznanie sztucznej inteligencji z perspektywy antropologicznej i etycznej, co miałoby służyć realizacji postulatu, że sztuczna inteligencja szanuje godność ludzką i promuje integralny rozwój osoby i społeczeństwa (a nie samą AI). Z kolei AI Act, z perspektywy prawnej, stara się o zapewnienie, aby systemy sztucznej inteligencji były bezpieczne, etycz-

ne i niezawodne oraz aby były wykorzystywane w sposób odpowiedzialny. Tak więc, mimo iż podejścia się różnią, wspólną płaszczyzną jest konstatacja czy imperatyw, że sztuczna inteligencja musi szanować podstawowe wartości ludzkości i stawiać człowieka w centrum.

W obu dokumentach zgodnie uznaje się potrzebę korzystania ze sztucznej inteligencji w sposób odpowiedzialny i w obrębie silnych ram etycznych. Jest to więc wyakcentowanie prawdy o etycznej dojrzałości ludzkich podmiotów, które pozostają w interakcji z AI i ponoszą odpowiedzialność za jej właściwe kształtowanie oraz użytkowanie. AeN kładzie przy tym duży nacisk na ludzką inteligencję w odróżnieniu od sztucznej, natomiast AI Act wymaga „człowieka w pętli”, czyli stałego nadzoru ze strony człowieka w istotnych i krytycznych punktach stosowania sztucznej inteligencji. Jako narzędzie do realizacji swoich postulatów AeN podkreśla znaczenie etyki (Dykasteria Nauki Wiary i Dykasteria ds. Kultury i Edukacji, 2025, nos. 36–48), w której chrześcijańskie zasady moralne kierują tworzeniem techniki, natomiast AI Act nakłada na systemy wysokiego ryzyka wymogi dotyczące przejrzystości, kontroli ze strony człowieka, a także audytu. W obydwu dokumentach kładzie się nacisk na zapobieganie dyskryminacji algorytmicznej i zapewnienie wytlumaczalności systemów sztucznej inteligencji. W przypadku dokumentu unijnego ustanowiono ponadto szczegółowe regulacje prawne prowadzące do osiągnięcia tego celu. Tak więc gdy Stolica Apostolska proponuje refleksję etyczną i filozoficzną opartą na godności ludzkiej i moralności chrześcijańskiej, UE ustanawia ramy regulacyjne z konkretnymi przepisami, aby zapewnić, że sztuczna inteligencja jest bezpieczna, przejrzysta i zgodna z prawami podstawowymi. Są to więc różne drogi, które mają doprowadzić do tego samego celu, jakim jest sztuczna inteligencja działająca zgodnie z zasadami afirmującymi osobę ludzką. Odmienność dróg widać także w różnych adresatach umocowania zobowiązań: podczas gdy Stolica Apostolska odwołuje się do jednostkowego sumienia i etyki indywidualnej, UE opowiada się za konkretną, w jakieś mierze przymusową regulacją, a więc za optyką społeczną.

Niezależnie od odmiennych spojrzeń na problem, obydwa dokumenty w centrum stawiają ostateczną odpowiedzialność człowieka na każdym etapie kontaktu z AI: w fazie jej tworzenia i definiowania, na etapie jej aplikacji do różnych obszarów życia i wykorzystywania indywidualnego oraz społecznego (przeciw dyskryminacji, marginalizacji), a także z perspektywy dalekosiężnych konsekwencji, które wiążą się z korzystaniem ze sztucznej inteligencji. Zamiast komentarza podsumowującego wagę wyzwania warto zacytować słowa Fidji Simo, nowej dyrektorki generalnej ds. aplikacji w OpenAI (zaliczona do 100 najważniejszych postaci AI w kategorii: przywódcy), która napisała: „Wierzę, że AI stwo-

rzy więcej możliwości dla większej liczby osób niż jakakolwiek inna technologia w historii. Jeśli wykonamy swoją pracę dobrze, AI może upodmiotowić wszystkich ludzi, w stopniu większym niż było to dotąd możliwe” (cyt. za: Duch, 2025). Nie powinien jednak ująć uwadze warunkowy charakter wypowiedzi, który uzależnia pozytywne skutki i konsekwencje stosowania AI od „dobrze wykonanego zadania” na poziomie etycznym i prawnym. Nie ulega wątpliwości, że na obecnym etapie rozwoju sztucznej inteligencji (przynajmniej w Europie) o to właśnie toczy się batalia.

4. Etyczne punkty zapalne, czyli perspektywy (i zagrożenia) na przyszłość

4.1. Paradoks samotności

Zazwyczaj listę zagrożeń otwierają kwestie społeczne, jak choćby zmiany na rynku pracy. Chciałbym jednak zaproponować wyjście od innego zjawiska, na które coraz częściej zwraca się uwagę, mianowicie od „paradoksu samotności”. W przestrzeni publicznej coraz częściej pojawiają się naukowo uzasadnione głosy, że z powszechnym „usieciowieniem” w ramach świata cyfrowego w żadnej mierze nie idzie w parze poszerzenie relacji i zmniejszenie alienacji jednostki. Wręcz przeciwnie, badania pokazują, iż możemy być „najbardziej samotnym pokoleniem” w historii ludzkości. Z wyników randomizowanych badań wynika, że szerokie korzystanie z chatbotów AI może pogłębić (lub pogłębia) poczucie samotności. Zdaniem naukowców, częste, codzienne korzystanie z chatbota koreluje z większym poczuciem samotności, z uzależnieniem emocjonalnym i z ograniczoną socjalizacją. Jeszcze bardziej do myślenia daje drugi wniosek z badań: że taki efekt generowany jest przez wszystkie rodzaje interakcji ze sztuczną inteligencją, a więc przez interakcje swobodne, emocjonalne i informacyjne. Poszukując schematu zachowań, badacze doszli do wniosku, że problematyczne wzorce używane na tym obszarze przypominają uzależnienie behawioralne i przemieszczenie emocjonalne (Fang *et al.*, 2025).

Komentarz moralny dotyczący tego zjawiska odnajdujemy już w encyklice papieża Franciszka z czasów pandemii „Fratelli tutti”. . (Franciszek, 2020, p. 43). Papież, demaskując mechanizmy stojące za relacjami wirtualnymi, wyjaśnia: „media cyfrowe mogą narazić na ryzyko uzależnienia, izolacji i postępującej utraty kontaktu z rzeczywistością, utrudniając rozwój autentycznych relacji międzyludzkich. Potrzeba gestów fizycznych, mimiki, milczenia, mowy ciała, a nawet zapachu, drżenia rąk, rumieńca, potu, ponieważ to wszyst-

ko mówi i należy do komunikacji międzyludzkiej. Kontakty wirtualne, które zwalniają ze żmudnego pielęgnowania przyjaźni, ze stabilnej wzajemności, a także z dojrzewającej z czasem zgodności, mają pozory kontaktów towarzyskich. Nie budują prawdziwie ‘nas’, ale zazwyczaj maskują i wzmacniają ten sam indywidualizm, który wyraża się w ksenofobii i pogardzie dla słabych. Połączenie cyfrowe nie wystarcza do budowania mostów, nie jest w stanie zjednoczyć ludzkości” (Franciszek, 2020, p. 43).

Drugie oblicze pseudokomunikacji cyfrowej to zawłaszczenie prawa do prywatności, które jest wartością osobową i etyczną: „W komunikacji cyfrowej dąży się do pokazania wszystkiego, a każda jednostka staje się obiektem spojrzeń, które rewidują, obnażają i rozpowszechniają, często anonimowo. Szacunek dla drugiego człowieka całkiem się rozpada, a w ten sposób właśnie, podczas gdy tego człowieka odsuwam, ignoruję i trzymam na dystans, mogę bezwstydnie wtargnąć w jego życie, aż do skrajności” (Franciszek, 2020, p. 42).

Ignorowanie tej prawdy i zaklinanie rzeczywistości przez przekonywanie o więziotwórczej roli komunikatorów wirtualnych wydaje się jednym z palących wyzwań, przed którymi stajemy, próbując wytyczyć granice dobra w świecie sztucznej inteligencji. Grupą szczególnie zagrożoną jest młode pokolenie (dzieci i młodzież), które naturalnie „wrosło” w świat wirtualny i traktuje sztuczną inteligencję jako „prawie ożywionego” partnera w komunikacji, nie rozumiejąc niebezpieczeństw, o których wspomniano wyżej.

4.2. Zmiany na rynku pracy

Dynamika zmian na rynku pracy ma swój wymierny komponent etyczny, jeśli wśród praw podstawowych wskażemy prawo człowieka do pracy. Od premiery ChatGPT w roku 2022 dają się słyszeć prognozy dotyczące turbulencji na rynku pracy w związku z narzędziami generatywnymi. Rozciągają się one od hiobowej wizji rychłego upadku rynku pracy po uspokajające przewidywania i wyniki badań innych podmiotów badawczych. Ilustracyjnie przywołajmy po jednym przedstawicielu obydwu opcji. W marcu 2023 r. firma Goldman Sachs – jedna z największych instytucji finansowych na świecie – ogłosiła szacunkowe prognozy dotyczące radykalnej transformacji systemów zawodowych, która miałaby nastąpić w ciągu około 10 lat (Kelly, 2023). Według tych szacunków około 300 mln stanowisk pracy zostanie zautomatyzowanych (jest to więcej niż liczy cała populacja osób w wieku produkcyjnym w Europie – ok. 260 mln w roku 2022). Jako możliwy scenariusz wskazano redukcję zatrudnienia, konieczność przekwalifikowania pracowników, krótszy tydzień pracy (aktualnie mowa jest o 4-dniowym) z ekonomicznymi konsekwencjami tych zjawisk. Okazało się jed-

nak, że kiedy doszło do wzmożonej redukcji etatów w dużych firmach typu Amazon czy YouTube, dyrektor generalny Goldman Sachs, nie widział powodów do paniki, a jedynie podkreślał bezprecedensowe tempo zmian na rynku pracy (Kubera, 2025).

Z drugiej strony, najnowsze badania przeprowadzone na Uniwersytecie Yale z października 2025 r. wskazują, że wpływ narzędzi generatywnych na rynek pracy jest znikomy (Gimbel, *et al.*, 2025). Na tym etapie dynamicznego rozwoju wykorzystania AI w różnych obszarach aktywności ludzkiej nie powinno dziwić, że prognozy dotyczące „nieuniknionych” rewolucyjnych zmian, formułowane są na przemian z uspokajającymi przewidywaniami różnych graczy na rynkach pracy i ekonomii. Istotne jest jednak, żeby monitorować dynamikę i owoce przemian generowanych z wykorzystaniem sztucznej inteligencji oraz łagodzić ewentualne negatywne skutki indywidualne i społeczne.

4.3. Życie człowieka w czasie działań zbrojnych

Brak poczucia bezpieczeństwa i niepokój dotyczący widma wojny w Europie, zwłaszcza w kontekście działań wojennych na terenie Ukrainy, każą pytać o rolę i znaczenie sztucznej inteligencji w nowoczesnej wojnie. Nie trzeba chyba przekonywać, że problem działań zbrojnych i prowadzonych wojen w dużej mierze narasta wraz z aktywnością AI w działaniach militarnych.

Z doniesień dziennikarskich wynika, że w wojnie na Ukrainie około 70 do 80% przypadków działań militarnych jest „produktem” AI (Sendek and Zalesiński, 2025), a AI-wspomagana nawigacja potrafi zwiększyć skuteczność uderzeń dronów z ~10–20% do ~70–80% (Bondar, 2025, p. 2).

Oznacza to, że zwłaszcza w atakach dronów maszyna sama decyduje, który z obiektów zaatakuje. Korzysta ona ze wpisanych algorytmów, ale „ostateczny wybór” obiektu pozostaje w gestii rachunku prawdopodobieństwa wyliczonego przez AI. W wymiarze etycznym formułowane są więc regulacje domagające się stałej kontroli człowieka nad działalnością sztucznej inteligencji w dziedzinie zbrojeń czy użycia broni, ponieważ „niezależność decyzyjna botów” jest ewidentnie groźna i śmiertelna, a wymóg ochrony ludności cywilnej, dzieci i osób najsłabszych w rachunku matematycznym staje się bezprzedmiotowy.

4.4. Animal rationale? Edukacyjne ostrzeżenie dla naszej „przyszłości”.

Wiele już napisano na temat możliwości wykorzystania sztucznej inteligencji w edukacji, w stworzeniu nowych możliwości, form i płaszczyzn kształcenia. Ich

adaptacja do procesu edukacyjnego może być szansą, którą należy wykorzystać dla dodatkowej stymulacji uczniów i studentów, a także dla inkluzji osób wykluczonych i z upośledzeniami. Tu również trzeba jednak pamiętać o ambiwalencji takiego zastosowania i pokazać drugą stronę medalu. Powaga nowych zagrożeń edukacyjnych mocno wybrzmiewa w poniższym głosie specjalisty: „Sztuczna inteligencja sprawi, że pisemne prace domowe stracą na znaczeniu. [...] Z perspektywy neuronauki wygląda to natomiast tak, że jeśli uczniowie i studenci zdecydują się na „outsourcing” własnego myślenia do zautomatyzowanych chatbotów, utracą umiejętność wyrażania tego, co myślą. Jeśli sztuczna inteligencja wyprze myślenie w placówkach oświatowych, może to w ciągu kilku dekad spowodować upadek naszej kultury opierającej się na istnieniu specjalistów, których mózgi były „trenowane” czy też kształcone poprzez szereg procesów wieloletniego uczenia się. [...] Ci, którym odmawia się szans na wszechstronną edukację, ostatecznie staną się ofiarami sztucznej inteligencji i jej konsekwencji społecznych” (Spitzer, 2025, pp. 240–241). Wydaje się, że takiego scenariusza należy się bardziej obawiać, niż wyobrażeń o futurystycznym buncie maszyn.

Jako pozytywny impuls dla przyszłości naszego życia ze sztuczną inteligencją warto zauważyć inicjatywę etyczną środowiska naukowców wiedeńskich, którzy opracowali 10 Zasad Moralnych dla Świata Cyfrowego (Spiekermann *et al.*, 2025) – na wzór biblijnego Dekalogu. Te zasady w dużej mierze diagnozują źródła etycznych wyzwań związanych z obecnością sztucznej inteligencji w naszej codzienności:

1. Nie podnoś technologii cyfrowej do rangi celu samego w sobie.
2. Nie przypisuj niesłusznie człowieczeństwa maszynom.
3. Stwórz przestrzeń na przestoje i spotkania analogowe.
4. Szanuj umiejętności społeczne i demokratyczne.
5. Nie niszczy natury dla postępu technologicznego.
6. Nie traktuj ludzi jak zwykłych obiektów danych.
7. Nie pozbawiaj siebie i innych ich prawdziwego ludzkiego potencjału.
8. Nie zaprzeczaj ograniczeniom technologii.
9. Nie podważaj wolności innych za pomocą środków technicznych.
10. Zapobiegaj koncentracji władzy i zapewnij uczestnictwo.

Odpowiedź na pytanie o etyczny wymiar funkcjonowania w świecie, w którym działa sztuczna inteligencja, polega na szukaniu rozwiązania drogą przybliżeń. Liczba niewiadomych, nieznane scenariusze rozwojowe, a szczególnie człowiek zmieniający się pod wpływem AI generują wiele możliwych rozwiązań i dróg dojścia. Niniejszy przyczynek jest także jednym z elementów tej etycznej układanki. Sam zabieg intelektualny przypomina w pewnej mierze doświadcze-

nie kuchni molekularnej lub witraża. Drobne impulsy pozwalają poznać smak; witraż – z drobnych fragmentów tworzy się obraz.

Data wpłynięcia: 2025-02-26;

Data uzyskania pozytywnych recenzji: 2025-12-08;

Data przesłania do druku: 2025-12-12.

References

- Bierzyński, T. (2025a) 'Ideological orientations of artificial intelligence towards traditional family models: A comparative study of twelve AI chatbots', *Wychowanie w Rodzinie*, 32(2), 171–188. Available at: <https://doi.org/10.61905/wwr/209801> (Accessed: December 3, 2025).
- Bierzyński, T. (2025b) *Strategie językowe zarządzania ograniczeniami w systemach sztucznej inteligencji stosowanych w cyfrowym poradnictwie rodzinnym: analiza porównawcza pięciu chatbotów AI*, 'Społeczeństwo – Kultura – Wartości. Studium społeczne', Nr 27, pp. 1–20. Available at: https://www.researchgate.net/publication/396431238_Strategie_jezykowe_zarzadzania_ograniczeniami_w_systemach_sztucznej_inteligencji_stosowanych_w_cyfrowym_poradnictwie_rodzinnym_analiza_porownawcza_pieciu_chatbotow_AI (Accessed: December 3, 2025).
- Bondar, K. (2025) 'Ukraine's Future Vision and Current Capabilities for Waging AI-Enabled Autonomous Warfare', *Center for Strategic & International Studies*, (03), pp. 1–37.
- Bugajski, J. (2023) 'Sztuczna inteligencja – zagrożenie czy bezpieczeństwo?', *Zbliżenia Cywilizacyjne*, 19(3), pp. 31–59. Available at: <https://doi.org/10.21784/ZC.2023.015>.
- Darmoros, K. (2025) *Papież o AI: mądrość to więcej niż dostęp do danych* – Vatican News. Available at: <https://www.vaticannews.va/pl/papiez/news/2025-06/papiez-o-ai-madrosco-to-wiecej-niz-dostep-do-danych.html> (Accessed: December 3, 2025).
- Duch, W. (2025) 'Pięć filarów rozwoju sztucznej inteligencji w Polsce,' *Wszystko co najważniejsze* [Pre-print]. Available at: <https://wszystkoconajwazniejsze.pl/prof-wlodzislaw-duch-sztuczna-inteligencja-w-polsce/> (Accessed: December 3, 2025).
- Dykasteria Nauki Wiary i Dykasteria ds. Kultury i Edukacji (2025) 'Antiqua et nova. Nota na temat relacji pomiędzy sztuczną inteligencją a inteligencją ludzką'. Available at: https://www.oaza.pl/wp-content/uploads/PL_Nota-Antiqua-et-nova.pdf (Accessed: December 3, 2025).
- eKAI (2025) 'Magazyn "Time": papież Leon XIV – myśliciel stawiający czoło rewolucji AI,' *eKAI | Portal Katolickiej Agencji Informacyjnej*, 2 September. Available at: <https://www.ekai.pl/magazyn-time-papiez-leon-xiv-mysliciel-stawiajacy-czolo-rewolucji-ai/> (Accessed: December 3, 2025).
- eKAI.pl (2025) *KAI – Serwis prasowy*. Available at: <https://serwiskai.pl/telegram/668309/> (Accessed: December 4, 2025).

- Eur-lex.europa.eu (2024) *Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/1689 z dnia 13 czerwca 2024 r. w sprawie ustanowienia zharmonizowanych przepisów dotyczących sztucznej inteligencji oraz zmiany rozporządzeń (WE) nr 300/2008, (UE) nr 167/2013, (UE) nr 168/2013, (UE) 2018/858, (UE) 2018/1139 i (UE) 2019/2144 oraz dyrektyw 2014/90/UE, (UE) 2016/797 i (UE) 2020/1828 (akt w sprawie sztucznej inteligencji) (Tekst mający znaczenie dla EOG)*. Available at: <http://data.europa.eu/eli/reg/2024/1689/oj> (Accessed: December 3, 2025).
- Fang, C.M. et al. (2025) *How AI and Human Behaviors Shape Psychosocial Effects of Chatbot Use: A Longitudinal Controlled Study*, MIT Media Lab. Available at: <https://www.media.mit.edu/publications/how-ai-and-human-behaviors-shape-psychosocial-effects-of-chatbot-use-a-longitudinal-controlled-study/> (Accessed: December 3, 2025).
- Franciszek (2020) *Encyklika Fratelli tutti Ojca Świętego Franciszka "O braterstwie i przyjaźni społecznej"*. Kraków: Wydawnictwo M.
- Gimbel, M. et al. (2025) *Evaluating the Impact of AI on the Labor Market: Current State of Affairs | The Budget Lab at Yale*. Available at: <https://budgetlab.yale.edu/research/evaluating-impact-ai-labor-market-current-state-affairs> (Accessed: December 4, 2025).
- Ilnicka, W. (2024) 'Sztuczna inteligencja – korzyści i zagrożenia,' *Zeszyty Naukowe Wyższej Szkoły Bankowej w Poznaniu*, 105(2). Available at: <https://doi.org/10.58683/dnswsb.2013>.
- Jacobs, S. (2025) *How We Chose the TIME100 AI 2025*, TIME. Available at: <https://time.com/7312067/how-we-chose-time100-ai-2025/> (Accessed: December 3, 2025).
- Kelly, J. (2023) *Goldman Sachs Predicts 300 Million Jobs Will Be Lost Or Degraded By Artificial Intelligence*, Forbes. Available at: <https://www.forbes.com/sites/jackkelly/2023/03/31/goldman-sachs-predicts-300-million-jobs-will-be-lost-or-degraded-by-artificial-intelligence/> (Accessed: December 3, 2025).
- Kuciński, M. (2021) 'Sztuczna inteligencja a godność człowieka,' *Fides, Ratio et Patria. Studia Toruńskie*, (14), pp. 210–224. Available at: <https://doi.org/10.56583/frp.650>.
- Kubera, G. (2025) *Szef Goldman Sachs nie kupuje paniki wokół miejsc pracy w erze AI*, Business Insider Polska. Available at: <https://businessinsider.com.pl/praca/szef-goldman-sachs-nie-kupuje-paniki-wokol-miejsc-pracy-w-erze-ai/z3k69hz> (Accessed: December 3, 2025).
- Lennox, J.C. (2023) *2084. Sztuczna inteligencja i przyszłość ludzkości*. Translated by Z. Kościuk. Warszawa: Fundacja Prodoteo (XX Pierwszy).
- Leo XIV (2025) *Message of the Holy Father to participants in the "Builders AI Forum 2025"*. Available at: <https://press.vatican.va/content/salastampa/en/bollettino/pubblico/2025/11/07/251107a.html> (Accessed: December 3, 2025).
- Rasińska, A. (2023) "Tomasz Rożek: próbowano nauczyć AI moralności, to zawsze wychodziło źle | eKAI,' *eKAI | Portal Katolickiej Agencji Informacyjnej*, 11 May. Available at: <https://www.ekai.pl/tomasz-rozek-probowano-nauczyc-ai-moralnosci-to-zawsze-wychodzilo-zle/> (Accessed: December 4, 2025).
- Salar, M.J. (2025) 'Antiqua et nova and the AI Regulation 2024: Tradition and future in the regulation of artificial intelligence,' *Bioethics Observatory – Institute of Life Sciences – UCV*, 25 February.

- Available at: <https://bioethicsobservatory.org/2025/02/antiqua-et-nova-and-the-ai-regulation-2024-tradition-and-future-in-the-regulation-of-artificial-intelligence/47592/> (Accessed: December 3, 2025).
- Sendek, R. and Zalesiński, Ł. (2025) *Polska Zbrojna*. Available at: <https://polska-zbrojna.pl/Mobile/ArticleShow/43398> (Accessed: December 3, 2025).
- Shah, S. (2025) *Pope Leo XIV: The 100 Most Influential Climate Leaders of 2025*, *TIME*. Available at: <https://time.com/collections/time-100-climate-2025/7326585/pope-leo-xiv/> (Accessed: December 4, 2025).
- Skalfist, P., Mikelsten, D. and Teigens, V. (2020) *Sztuczna inteligencja: czwarta rewolucja przemysłowa*. S.l.: Cambridge Stanford Books. Available at: <https://ebook.yourcloudlibrary.com/library/oclc/detail/3mbpwz9> (Accessed: December 4, 2025).
- Ślipko, T. (1974) *Etos chrześcijański. Zarys etyki ogólnej*. Kraków: Wydawnictwo Apostolstwa Modlitwy.
- Spiekermann, S. et al. (2025) 'The Power of 10: New Rules for the Digital World,' in. *The Human Person nad their IA*, Ljubana. Available at: <https://doi.org/10.13140/RG.2.2.28008.38403>.
- Spitzer, M. (2025) *Sztuczna Inteligencja. Ponad człowiekiem. AI jako ratunek i zagrożenie*. Warszawa: Dobra Literatura (Kontrasty i kontrowersje).
- Sypniewska, B.A. and Gołębiowski, G. (2023) "Sztuczna inteligencja – dylematy etyczne,' *Przegląd Organizacji*, pp. 248–254. Available at: <https://doi.org/10.33141/po.2023.03.26>.
- Szpyrka, M. (2025) 'Halucynacje sztucznej inteligencji: wyzwania, przyczyny i przeciwdziałanie dezinformacji,' *EURACTIV.pl*, 28 January. Available at: <https://euractiv.pl/section/nowe-technologie/news/halucynacje-sztucznej-inteligencji-wyzwania-przyczyny-i-przeciwdzialanie-dezinformacji/> (Accessed: December 3, 2025).
- Wiara.pl (2025) *Leon XIV: AI powinna odzwierciedlać Boży zamysł*, *Wiara.pl Kultura*. Available at: <https://kultura.wiara.pl/doc/9486518.Leon-XIV-AI-powinna-odzwierciedlac-Bozy-zamysl> (Accessed: December 4, 2025).
- Wyrostkiewicz, M. (2025) 'Dezinformacja w kontekście troski o pokój: Studium na marginesie papieskiego Orędzia na LVIII Światowy Dzień Pokoju (2025),' *Collectanea Theologica*, 95(2), pp. 359–389. Available at: <https://doi.org/10.21697/ct.2025.95.2.05>.