

Twój przewodnik po

VALUES -

Umiejętności AI
wspierane przez
VALUES



www.valuesai.eu

2025
Studia
przypadków

Twórca
Generation (Change?)



Co-funded by
the European Union

spis treści

01	Powitanie
04	Projekt
06	Studia przypadków
16	Źródła
18	Wnioski
19	Kontakt



Co-funded by
the European Union

Finansowane przez Unię Europejską. Wyrażone poglądy i opinie są jednak poglądami i opiniami wyłącznie autora (autorów) i niekoniecznie odzwierciedlają poglądy Unii Europejskiej ani PGeneration (Change?). Ani Unia Europejska, ani Generation (Change?) nie ponoszą za nie odpowiedzialności.

01

VALUES: przewodnik po studiach przypadków



Projekt

VALUES to projekt Erasmus+, którego celem jest wyposażenie młodych Europejczyków w wiedzę, umiejętności i zrozumienie etyki niezbędne do odniesienia sukcesu w przyszłości kształtowanej przez sztuczną inteligencję. W obliczu transformacji sztucznej inteligencji, która zmienia sposób, w jaki pracujemy, uczymy się i żyjemy, VALUES dba o to, aby młodzi ludzie byli nie tylko uczestnikami, ale i liderami tej cyfrowej rewolucji.

Dzięki najnowocześniejszym szkoleniom, innowacyjnym narzędziom i kompleksowemu Pakietowi Startowemu dla Etycznej AI, łączymy technologię z ludzkością. Nasza interaktywna Platforma Samooceny Umiejętności AI oraz dynamiczne Otwarte Zasoby Edukacyjne zostały zaprojektowane tak, aby nauka była dostępna, angażująca i skuteczna.

VALUES nie dotyczy wyłącznie edukacji — chodzi o kształtowanie pokolenia odpowiedzialnych, kierujących się wartościami obywateli cyfrowych, którzy potrafią etycznie wykorzystać sztuczną inteligencję do stawiania czoła wyzwaniom społecznym, budowania inkluzywnych społeczności i opracowywania zrównoważonych rozwiązań.

Partnerzy projektu:

- Uniwersytet Szczeciński (USA) koordynator projektu
- Stowarzyszenie Wspierania Techniki Polskiej (SWTP)
- Generation (Change?)
- Momentum Marketing Services (MMS)
- European E-Learning Institute (EUEI)
- The Vision Works (tw GmbH)



Odpowiedzialna sztuczna
inteligencja nie polega
tylko na
odpowiedzialności –
chodzi o zapewnienie, że
to, co tworzysz,
umożliwia rozwój
człowieka



Rumman Chowdhury,
Dyrektor generalny Parity AI, 2023





Studium przypadku 1

Skandal algorytmiczny na egzaminach A-level w Wielkiej Brytanii: lekcja o odpowiedzialnej sztucznej inteligencji w edukacji

Streszczenie: W 2020 roku, z powodu odwołania egzaminów A-level z powodu pandemii COVID-19, rząd Wielkiej Brytanii wdrożył algorytm przyznawania ocen uczniom. Algorytm ten, oparty na historycznych wynikach szkół i danych rankingowych, miał na celu ujednolicenie wyników. Szybko jednak stało się jasne, że system faworyzował uczniów z dobrze prosperujących, często zamożniejszych szkół, kosztem tych z mniej uprzywilejowanych środowisk. Wielu uczniów otrzymało oceny znacznie niższe niż te przewidywane przez ich nauczycieli, co wywołało ogólnokrajowe protesty i oburzenie społeczne. Ostatecznie rząd zmienił swoją decyzję i przywrócił oceny przewidywane przez nauczycieli.

Najważniejsze wnioski:

- **Przejrzystość i możliwość audytu:** Systemy sztucznej inteligencji w edukacji muszą być przejrzyste i regularnie poddawane audytom, aby zapewnić ich uczciwość.
- **Projektowanie uwzględniające kontekst:** Algorytmy powinny być opracowywane z uwzględnieniem różnorodności społecznej i nie mogą wzmacniać istniejących uprzedzeń.

Temat studium przypadku:

- Egzaminy o wysokiej stawce (A-levels) w Wielkiej Brytanii
- Nierówności społeczne i stronniczość algorytmiczna
- Wdrażanie odpowiedzialnej sztucznej inteligencji (RAI)
- Etyka i przejrzystość w algorytmach edukacyjnych

Znaczenie dla czytelnika: To studium przypadku podkreśla kluczowe znaczenie odpowiedzialnego projektowania i wdrażania systemów AI w edukacji. Ilustruje ono, jak brak przejrzystości, pomijanie kontekstu społecznego i nadmierne poleganie na danych historycznych mogą pogłębiać istniejące nierówności. Dla decydentów, nauczycieli i twórców technologii edukacyjnych stanowi to przestrożę przed bezkrytycznym stosowaniem algorytmów w procesach decyzyjnych, zwłaszcza w tak wrażliwych obszarach, jak ocena uczniów.

Najważniejsze wnioski:

- **Zaangażowanie interesariuszy:** Nauczyciele, uczniowie i społeczności powinni brać udział w projektowaniu i wdrażaniu systemów AI.
- **Reaktywność i odpowiedzialność:** Instytucje muszą być gotowe na szybką reakcję na niezamierzone konsekwencje decyzji algorytmicznych i szybką ich korektę.

Studium przypadku 2

Stronniczy algorytm rekrutacyjny Amazon – lekcja o stronniczości sztucznej inteligencji

Streszczenie: W 2014 roku zespół specjalistów ds. uczenia maszynowego w firmie Amazon rozpoczął prace nad innowacyjnym narzędziem rekrutacyjnym opartym na sztucznej inteligencji. Celem projektu była automatyzacja procesu selekcji CV i identyfikacja najlepszych kandydatów na stanowiska techniczne, w tym inżynierów oprogramowania. System oceniał aplikacje w skali od jednej do pięciu gwiazdek – podobnie jak klienci oceniają produkty na platformie Amazon. Do 2015 roku zespół odkrył poważny problem: algorytm był uprzedzony w stosunku do kobiet aplikujących na stanowiska techniczne. Narzędzie zostało przeszkolone na podstawie życiorysów przesłanych do firmy w ciągu poprzedniej dekady – z których większość pochodziła od mężczyzn, co odzwierciedlało dominację mężczyzn w branży technologicznej. W rezultacie algorytm „nauczył się”, że kandydaci płci męskiej są preferowani. System karał życiorysy zawierające słowo „kobiecy” (np. „kapitan kobiecego klubu szachowego”) i obniżał oceny absolwentek dwóch uczelni wyłącznie dla kobiet. Faworyzował również kandydatów, którzy w swoich CV używali czasowników „męskich”, takich jak „wykonał” i „złapał”. Pomimo prób inżynierów Amazona, aby wyeliminować ten błąd, firma ostatecznie porzuciła projekt w 2018 roku, przyznając, że nie ma gwarancji, że system nie rozwinie innych form dyskryminacji. Sprawa ta stała się kluczową lekcją dla branży technologicznej na temat ryzyka związanego z uprzedzeniami algorytmicznymi i znaczenia odpowiedzialnego rozwoju sztucznej inteligencji.

Temat studium przypadku:

- Sztuczna inteligencja w rekrutacji
- Uprzedzenia algorytmu
- Etyka sztucznej inteligencji
- Równość płci w technologii

Znaczenie dla czytelnika: Algorytm rekrutacyjny Amazona stanowi kluczowe studium przypadku dotyczącego stronniczości algorytmicznej w systemach sztucznej inteligencji. Pokazuje on, jak systemy uczenia maszynowego mogą nie tylko odzwierciedlać, ale także wzmacniać istniejące nierówności społeczne, gdy są trenowane na stronniczych danych historycznych. W czasach, gdy sztuczna inteligencja jest coraz częściej wykorzystywana w procesach rekrutacyjnych (55% specjalistów HR w USA przewiduje, że sztuczna inteligencja stanie się rutynowym elementem ich pracy w ciągu pięciu lat), zrozumienie potencjalnego ryzyka związanego z uprzedzeniami algorytmicznymi jest kluczowe zarówno dla pracodawców, jak i osób poszukujących pracy. Przypadek Amazon pokazuje, że nawet najbardziej zaawansowane technologicznie organizacje mogą nieumyślnie tworzyć dyskryminujące systemy. Studium przypadku podkreśla znaczenie różnorodności i integracji w rozwoju sztucznej inteligencji. Ilustruje ono, jak brak różnorodności w zespołach technologicznych i danych szkoleniowych może prowadzić do powstania systemów faworyzujących jedną grupę demograficzną kosztem innych. Dla specjalistów zajmujących się etyką sztucznej inteligencji, ten przypadek stanowi konkretny przykład wyzwań związanych z zapewnieniem uczciwości algorytmicznej. Podkreśla potrzebę rygorystycznego testowania systemów sztucznej inteligencji pod kątem potencjalnych uprzedzeń przed wdrożeniem, zwłaszcza w kontekstach wysokiego ryzyka, takich jak zatrudnienie. Ta historia jest również istotna w kontekście nowych przepisów dotyczących sztucznej inteligencji, takich jak nowojorska ustawa antydyskryminacyjna dotycząca algorytmów rekrutacyjnych, która weszła w życie w 2023 roku. Podkreśla ona potrzebę ustanowienia standardów i regulacji zapewniających uczciwość, przejrzystość i rozliczalność zautomatyzowanych systemów podejmowania decyzji.



Studium przypadku 3

Integralność i uczciwość egzaminów oparta na sztucznej inteligencji: platforma YouTestMe jako studium przypadku

Streszczenie: W tym studium przypadku przyjrzymy się platformie YouTestMe opartej na sztucznej inteligencji, która rozwiązuje problemy związane z uczciwością i rzetelnością egzaminów online.

Do najważniejszych innowacji należą:

- Automatyczne ocenianie za pomocą sztucznej inteligencji: automatyczne ocenianie prac pisemnych przy użyciu przetwarzania języka naturalnego (NLP) w celu oceny spójności semantycznej, integralności strukturalnej i poprawności językowej, co zmniejsza obciążenie egzaminatorów, a jednocześnie zapewnia spójność.
- Zautomatyzowany nadzór: sztuczna inteligencja monitoruje osoby zdające egzamin pod kątem podejrzanych zachowań (np. przełączanie okien, niezgodność twarzy) z dokładnością wykrywania wynoszącą 98%, co gwarantuje bezpieczeństwo egzaminu.
- Percepcje uczciwości: Badania pokazują, że uczniowie postrzegają algorytmy sztucznej inteligencji jako bardziej sprawiedliwe niż oceniający ludzie, szczególnie w przypadku ocen kształtujących
- Przejrzystość i dalsze wyjaśnienia po ocenie zwiększa zaufanie do wyników uzyskiwanych dzięki sztucznej inteligencji. Platforma jest zgodna z zasadami odpowiedzialnej sztucznej inteligencji, priorytetowo traktując prywatność danych, możliwość audytu i nadzór ludzki. Na przykład nagrania z nadzoru są przechowywane w celu opcjonalnej weryfikacji przez człowieka, łącząc wydajność sztucznej inteligencji z rozliczalnością.

Temat studium przypadku:

- Automatyczne ocenianie i nadzór egzaminów w trybie online
- Sprawiedliwość sztucznej inteligencji i ograniczanie uprzedzeń w ocenie edukacyjnej
- Współpraca między człowiekiem i AI na rzecz transparentnej oceny
- Odpowiedzialne wdrażanie sztucznej inteligencji w testowaniu o dużej wadze

Znaczenie dla czytelnika:

Edukatorzy/Instytucje: Pokazuje, jak sztuczna inteligencja może usprawnić procesy oceny, jednocześnie eliminując uprzedzenia. Badanie podkreśla, że postrzeganie sprawiedliwości sztucznej inteligencji poprawia się, gdy kryteria oceny są przejrzyste i zawierają wyjaśnienia.

Twórcy sztucznej inteligencji: Oferuje model umożliwiający integrację zabezpieczeń etycznych (np. sprawdzanie stronniczości danych szkoleniowych, protokoły zgody użytkownika) z systemami EdTech, jak widać na przykładzie zgodnego z RODO przetwarzania danych przez YouTestMe.

Decydenci: Ilustruje znaczenie standardów nadzoru AI i automatycznego oceniania, podkreślając potrzebę nadzoru ze strony człowieka w celu ograniczenia ryzyka, takiego jak hegemonia algorytmów.

Studium przypadku przedstawia praktyczne wskazówki dotyczące odpowiedzialnego wdrażania sztucznej inteligencji w ocenach, równoważąc efektywność z równością.



Studium przypadku 4

Narodowe Zgromadzenie Młodzieży na temat Sztucznej Inteligencji – Perspektywy Młodzieży w Polityce AI

Streszczenie:

W 2023 roku Narodowe Zgromadzenie Młodzieży ds. Sztucznej Inteligencji zwołało młodych ludzi w wieku od 12 do 24 lat, aby omówić irlandzką Narodową Strategię AI („AI – Here for Good”). Celem inicjatywy było zapewnienie uwzględnienia głosu młodzieży w kształtowaniu zarządzania, polityki i edukacji w zakresie AI.

Kluczowe obszary dyskusji obejmowały:

- Równość i integracja w sztucznej inteligencji – obawy dotyczące stronniczości i uczciwości w zastosowaniach sztucznej inteligencji.
- Wiarygodność sztucznej inteligencji – znaczenie przejrzystości i odpowiedzialności w procesie podejmowania decyzji przez sztuczną inteligencję.
- Regulacja sztucznej inteligencji i świadomość społeczna – potrzeba jasnych przepisów dotyczących wykorzystania sztucznej inteligencji w szkołach, miejscach pracy i mediach społecznościowych.

Uczestnicy zaproponowali rekomendacje dotyczące zapewnienia, aby sztuczna inteligencja wspierała rozwój młodzieży, integrację społeczną i sprawiedliwą reprezentację. Ich opinie zostały udokumentowane w raporcie, który posłużył jako podstawa do dyskusji na temat polityki rządowej.

Temat studium przypadku:

- Etyka i regulacje dotyczące sztucznej inteligencji
- Zaangażowanie młodzieży w politykę
- Świadomość i umiejętność posługiwania się sztuczną inteligencją

Znaczenie dla czytelnika:

Polityki dotyczące sztucznej inteligencji mają bezpośredni wpływ na młodzież, ale młodzi ludzie są często pomijani w procesie podejmowania decyzji.

Studium przypadku podkreśla znaczenie zaangażowania młodzieży w kształtowanie etyki i regulacji dotyczących sztucznej inteligencji.

Zachęca nauczycieli i pracowników młodzieżowych do propagowania wiedzy na temat sztucznej inteligencji, aby zapewnić, że młodzi ludzie rozumieją swoje prawa i obowiązki w świecie zdominowanym przez sztuczną inteligencję.

Pokazuje, w jaki sposób głosy młodych ludzi mogą przyczynić się do kształtowania polityki w zakresie sztucznej inteligencji — jest to praktyka, którą należy rozszerzyć na całą Europę.

Studium przypadku 5

Interwencje przeciwko cyberprzemocy oparte na sztucznej inteligencji – ocena perspektyw młodzieży

Streszczenie:

W niniejszym badaniu analizowano postrzeganie przez młodzież interwencji przeciwko cyberprzemocy w mediach społecznościowych, opartych na sztucznej inteligencji. Naukowcy z Dublin City University (DCU) przetestowali oparte na sztucznej inteligencji proaktywne strategie moderacji treści, mające na celu wykrywanie i ograniczanie szkodliwych interakcji online. W konsultacjach wzięli udział młodzi ludzie w wieku od 12 do 17 lat, którzy ewaluowali te interwencje za pośrednictwem grup fokusowych i dyskusji online.

Badanie wykazało, że chociaż moderacja treści przez sztuczną inteligencję może ograniczyć szkodliwe interakcje, młodzi uczestnicy wyrazili obawy dotyczące:

- Fałszywe alarmy i cenzura – sztuczna inteligencja nieprawidłowo sygnalizuje nieszkodliwe treści jako szkodliwe.
- Brak kontekstu ludzkiego – sztuczna inteligencja ma problemy ze zrozumieniem niuansów humoru i języka nieformalnego.
- Obawy dotyczące prywatności – niepewność co do tego, kto kontroluje interwencje sztucznej inteligencji i ich wpływ na prawa cyfrowe.

Badania podkreśliły potrzebę udziału młodzieży w projektowaniu sztucznej inteligencji, tak aby moderacja treści była zgodna z perspektywą młodych ludzi i cyfrową rzeczywistością.

Temat studium przypadku:

- Bezpieczeństwo w Internecie
- Stronniczość w moderacji AI
- Cyfrowe dobre samopoczucie młodzieży

Znaczenie dla czytelnika:

Wielu młodych ludzi korzysta z platform opartych na sztucznej inteligencji w celu interakcji społecznych i edukacji, co sprawia, że etyczne wdrażanie sztucznej inteligencji ma kluczowe znaczenie.

Cyberprzemoc to coraz poważniejszy problem i choć moderacja AI może pomóc, musi być transparentna, sprawiedliwa i rozliczalna wobec użytkowników. To studium przypadku zachęca pracowników młodzieżowych i edukatorów do dyskusji na temat roli AI w bezpieczeństwie cyfrowym, promując krytyczne myślenie o algorytmicznym podejmowaniu decyzji. • Studium podkreśla znaczenie angażowania młodych ludzi w projektowanie i ocenę interwencji AI, zapewniając, że rozwiązania są zorientowane na młodzież i etycznie uzasadnione.



Studium przypadku 6

Sztuczna inteligencja odpowiedzialna za równość i jakość testów: test z języka angielskiego Duolingo jako studium przypadku

Streszczenie:

Niniejszy artykuł, który ukaże się wkrótce w ramach Handbook for Assessment in the Service of Learning, wykorzystuje test języka angielskiego Duolingo (DET) jako studium przypadku, aby zbadać kluczową rolę odpowiedzialnej sztucznej inteligencji (RAI) w ocenie o wysokiej stawce. Chociaż sztuczna inteligencja (AI) oferuje wydajność w generowaniu i ocenianiu zadań, wiąże się również z ryzykiem, takim jak stronniczość. DET uwzględnia te zagrożenia poprzez cztery konkretne standardy RAI – Trafność i rzetelność, Uczciwość, Prywatność i bezpieczeństwo oraz Rozliczalność i przejrzystość – pokazując, jak przekładają się one na praktyczną implementację.

W artykule szczegółowo opisano rozwój tych standardów i ich związek z szerszymi zasadami RAI, podając konkretne przykłady tego, jak praktyki te zapewniają zarówno jakość testów (prawidłowe wnioski dotyczące wyników), jak i sprawiedliwość (sprawiedliwość dla wszystkich zdających). DET stanowi model tego, jak firmy mogą transparentnie i odpowiedzialnie integrować sztuczną inteligencję, jednocześnie minimalizując ryzyko etyczne.

Temat studium przypadku:

- Test znajomości języka angielskiego
- Sprawiedliwość i bezstronność sztucznej inteligencji w testach
- (Human in the loop) HiTL i AI
- Odpowiedzialna sztuczna inteligencja

Znaczenie dla czytelnika:

Znaczenie niniejszego artykułu dla czytelnika koncentruje się na kluczowej roli odpowiedzialnej sztucznej inteligencji (RAI) w zapewnianiu zarówno jakości, jak i równości ocen wspomaganych przez sztuczną inteligencję. Wykorzystuje on test języka angielskiego Duolingo (DET) jako przekonujące studium przypadku, szczegółowo opisując, w jaki sposób zasady RAI są wdrażane w praktyce, aby minimalizować ryzyko związane ze sztuczną inteligencją, jednocześnie maksymalizując jej korzyści. Jest to cenne dla każdego, kto zajmuje się tworzeniem ocen, etyką sztucznej inteligencji lub wykorzystaniem sztucznej inteligencji w testach wysokiego ryzyka. Artykuł zawiera szczegółowy, praktyczny przykład wdrożenia zasad RAI w rzeczywistych warunkach o dużym wpływie. Jest to kluczowe, biorąc pod uwagę rosnące wykorzystanie sztucznej inteligencji w edukacji i ocenie. Oferuje również wgląd w proces tworzenia oceny wspomaganej przez sztuczną inteligencję, w tym generowanie elementów, punktację i środki bezpieczeństwa. Podkreśla znaczenie podejścia z udziałem człowieka w utrzymaniu jakości i równości. Etyka i zarządzanie w dziedzinie sztucznej inteligencji (AI): Cztery przedstawione standardy RAI (ważność i niezawodność, uczciwość, prywatność i bezpieczeństwo, rozliczalność i przejrzystość) stanowią wzór dla organizacji opracowujących i wdrażających systemy AI, szczególnie w zastosowaniach o wysokim ryzyku. Zgodność z wytycznymi NIST AI RMF dodatkowo wzmacnia ważność tych ram. Artykuł pokazuje, jak systematyczne stosowanie praktyk RAI może przyczynić się do osiągnięcia zarówno równości w teście (sprawiedliwości dla wszystkich zdających), jak i jakości (trafności wniosków z wyników). Ma to kluczowe znaczenie dla zapewnienia, że oceny oparte na sztucznej inteligencji nie będą stronnicze ani dyskryminujące.



Studium przypadku 7

Dyskryminacja ze względu na płeć w reklamach na popularnych portalach społecznościowych.

Streszczenie:

Global Witness zainicjował płatne reklamy na Facebooku w następujących krajach: Wielkiej Brytanii, Holandii, Francji, Indiach, Irlandii i Republice Południowej Afryki. Reklamy koncentrowały się na prawdziwych ofertach pracy z różnych branż, w tym dla pilotów, fryzjerów i psychologów. Wyniki wyszukiwania reklam kierowanych przez algorytmy w mediach społecznościowych wykazały wyraźne uprzedzenia ze względu na płeć w kontekście wakacji: 91% osób, którym wyświetlono ogłoszenie o pracę mechanika, to mężczyźni, a 79% osób, którym wyświetlono ogłoszenie o pracę nauczyciela przedszkolnego, to kobiety.

Dalsze badania przeprowadzone w Holandii i Francji we współpracy z organizacją feministyczną wykazały, że 97% osób, którym pokazano ogłoszenie o pracę na stanowisku recepcjonisty, stanowiły kobiety.

W ramach obowiązkowego procesu selekcji podczas tworzenia reklamy firma Global Witness wybrała cel polegający na kierowaniu reklam do osób, które według Facebooka najprawdopodobniej klikną reklamę, przy czym przy każdej reklamie określono, że musi być ona wyświetlana osobom dorosłym mieszkającym w danym kraju lub niedawno przebywającym w jednym z tych krajów.

Temat studium przypadku:

- Media społecznościowe
- Stronniczość

Znaczenie dla czytelnika:

To studium przypadku uwypukla głęboko zakorzenione uprzedzenia obecne w algorytmach, które wpływają na treści, jakie ludzie widzą codziennie, zwłaszcza w ogłoszeniach o pracę w mediach społecznościowych. Systemy oparte na sztucznej inteligencji, które decydują o umiejscowieniu reklam, często odzwierciedlają i wzmacniają istniejące uprzedzenia społeczne, co prowadzi do rozbieżności w sposobie, w jaki oferty pracy są prezentowane różnym płciom.

Media społecznościowe wyewoluowały poza platformę do luźnej interakcji; obecnie stanowią główne źródło informacji, wsparcia i poradnictwa zawodowego. Jednak uprzedzenia algorytmiczne oznaczają, że niektóre stanowiska pracy mogą być nieproporcjonalnie częściej przydzielane mężczyznom lub kobietom na podstawie danych historycznych, a nie indywidualnych zasług, co ogranicza dostęp do możliwości i wzmacnia stereotypy.

Biorąc pod uwagę zależność młodych ludzi od mediów społecznościowych w zakresie informacji związanych z karierą zawodową, kluczowe jest podnoszenie świadomości na temat tego uprzedzenia i wyposażenie ich w umiejętności krytycznego myślenia, aby móc je rozpoznawać i przeciwdziałać. Zachęcanie do przejrzystości systemów sztucznej inteligencji i propagowanie bardziej sprawiedliwych praktyk algorytmicznych może pomóc w stworzeniu bardziej sprawiedliwego cyfrowego rynku pracy.



Studium przypadku 8:

Mówienie po niemiecku w wirtualnej rzeczywistości – wspieranie młodych uczniów języków za pomocą sztucznej

Streszczenie:

W ramach niedawnego niemieckiego projektu pilotażowego przetestowano wykorzystanie środowiska edukacyjnego opartego na wirtualnej rzeczywistości (VR) do wspierania dzieci i młodzieży uczących się języka niemieckiego jako drugiego języka (DaZ). Immersyjny scenariusz VR pozwolił uczniom ćwiczyć realistyczne, codzienne rozmowy w bezpiecznej i angażującej przestrzeni cyfrowej. Projekt był częścią szerszej inicjatywy mającej na celu ocenę, w jaki sposób narzędzia VR wspomagane sztuczną inteligencją i oparte na mowie mogą poprawić płynność, motywację i pewność siebie wśród młodzieży z ograniczoną znajomością języka niemieckiego.

Uczniowie w zestawach VR poruszali się po wirtualnym mieszkaniu, wchodząc w interakcje z awatarami AI, wykonując domowe obowiązki i odpowiadając na pytania. System zapewniał natychmiastową informację zwrotną i adaptacyjne wsparcie. Nauczyciele zaobserwowali poprawę płynności mówienia, szczególnie wśród uczniów, którzy zazwyczaj byli nieśmiali w klasach. Badanie wskazało jednak również na wyzwania: ograniczony dostęp do sprzętu VR, potrzebę szkoleń dla nauczycieli oraz obawy dotyczące inkluzji.

Temat studium przypadku:

- AI + VR w edukacji językowej
- Technologia immersyjna wspierająca zaangażowanie młodzieży i naukę języka w klasach wielojęzycznych
- Odpowiedzialne korzystanie z nowych technologii w edukacji

Znaczenie dla czytelnika:

Ten przypadek ilustruje, jak sztuczna inteligencja i technologie immersyjne mogą poprawić zaangażowanie młodzieży i wyniki w nauce języków obcych. Pokazuje potencjał nowych narzędzi do budowania pewności siebie, zachęcania do aktywnego używania nowych języków i zmniejszania lęku w klasie – zwłaszcza wśród uczniów migrantów lub nowo przybyłych.

Dla nauczycieli stawia to pytania praktyczne i etyczne: Kto ma dostęp do tych narzędzi? Czy interwencje oparte na VR są inkluzywne? Jakiego wsparcia potrzebują nauczyciele, aby odpowiedzialnie korzystać z takich narzędzi? Przypadek ten zachęca nauczycieli, twórców EdTech i rzeczników młodzieży do eksploracji VR jako motywującej przestrzeni edukacyjnej, przy jednoczesnym krytycznym spojrzeniu na dostępność i użyteczność.

Studium przypadku 9:

Strategia AfD na rzecz cyfrowej młodzieży – sztuczna inteligencja, TikTok i wzrost skrajnej prawicy

Streszczenie:

W ostatnich latach znacząco wzrosła atrakcyjność Alternatywy dla Niemiec (AfD) wśród młodych wyborców, szczególnie dzięki strategicznemu wykorzystaniu platform cyfrowych i treści generowanych przez sztuczną inteligencję. Młodzieżowe skrzydło partii, Młoda Alternatywa (JA), wykorzystywało narzędzia sztucznej inteligencji do tworzenia angażujących treści, w tym teledysków i kampanii w mediach społecznościowych, które znalazły oddźwięk wśród młodszych grup demograficznych. Pomimo uznania AfD za organizację ekstremistyczną przez Federalny Urząd Ochrony Konstytucji (BfV) w 2023 roku, strategię cyfrową JA przyczyniły się do znacznego wzrostu poparcia młodzieży dla AfD. Niniejsze studium przypadku analizuje metody stosowane przez AfD i JA w celu angażowania młodych wyborców, implikacje ich strategii cyfrowych oraz szerszy wpływ na niemiecką politykę.

Temat studium przypadku:

Punkt styku sztucznej inteligencji, mediów społecznościowych i radykalizacji politycznej wśród młodzieży.

Znaczenie dla czytelnika:

Ta sprawa uwypukla potencjał sztucznej inteligencji i platform cyfrowych w zakresie wpływania na poglądy i zachowania polityczne młodych ludzi. Podkreśla ona wagę edukacji medialnej i krytycznego podejścia do treści online, zwłaszcza w kontekście przekazów politycznych.

Zrozumienie tej dynamiki jest kluczowe dla edukatorów, decydentów i działaczy na rzecz młodzieży, którzy dążą do wspierania świadomego i demokratycznego uczestnictwa młodych obywateli.

- Strategiczne wykorzystanie sztucznej inteligencji i mediów społecznościowych:

AfD i jej młodzieżówka skutecznie wykorzystały treści generowane przez sztuczną inteligencję i platformy takie jak TikTok, aby tworzyć atrakcyjne przekazy skierowane do młodych wyborców. Na przykład JA wydała utwór taneczny zatytułowany „Remigration Hit” z towarzyszącym mu filmem przedstawiającym masowe deportacje w stylizowany, uroczysty sposób.

- Sklasyfikowanie jako organizacja ekstremistyczna:

W 2023 roku BfV uznało JA za pravicową grupę ekstremistyczną, powołując się na jej etnicznie zdefiniowaną koncepcję ludu i antyimigranckie stanowisko.

- Działania mające na celu rozwiązanie i zmianę marki:

W obliczu potencjalnych konsekwencji prawnych i publicznej kontroli, AfD rozwiązała JA na początku 2025 roku, dążąc do zdystansowania się od ekstremistów. Ogłoszono jednak plany utworzenia nowej organizacji młodzieżowej pod bezpośrednią kontrolą partii.



Studium przypadku 10:

Kiedy sztuczna inteligencja źle ocenia pracę domową – asystent Fobizz w klasie

Streszczenie:

Pod koniec 2023 roku niemiecka platforma technologii edukacyjnych Fobizz zaprezentowała wersję beta swojego „Asystenta Korekty AI” – narzędzia, którego celem jest pomoc nauczycielom w ocenie prac domowych i pisemnych w otwartym tekście z wykorzystaniem generatywnej sztucznej inteligencji. Promowane jako narzędzie wspomagające oszczędzanie czasu, szybko zyskało zainteresowanie nauczycieli eksperymentujących z nowymi formami automatyzacji zajęć w klasie.

Jednak badanie przeprowadzone przez niemieckich specjalistów ds. technologii edukacyjnych, o którym informowały liczne media na początku 2024 roku, ujawniło krytyczne błędy: asystent przyznawał wysokie oceny odpowiedziom bezsensownym lub nieistotnym, często nie wykrywał oczywistych błędów merytorycznych i wykazywał niespójności w ocenie podobnych odpowiedzi. W jednym z przypadków testowych przesłano nawet tekst zastępczy, taki jak „Lorem ipsum”, który otrzymał od narzędzia ocenę „dobrą”.

Chociaż narzędzie nie zostało oficjalnie wdrożone w systemach szkolnych, a jego twórcy potwierdzili jego status wersji beta, sprawa ta uwypukla poważne wyzwania związane z wdrażaniem rozwiązań do

oceniania opartych na sztucznej inteligencji, zwłaszcza w kontekście edukacji młodzieży i ocen o wysokiej stawce.

Temat studium przypadku:

- Sztuczna inteligencja generatywna w szkołach
- Błędna ocena AI i niezamierzone konsekwencje
- Odpowiedzialne testowanie AI w edukacji
- Zaufanie i nadzór we współpracy człowieka ze sztuczną inteligencją

Znaczenie dla czytelnika:

Ta sprawa jest szczególnie istotna dla nauczycieli, twórców technologii edukacyjnych i decydentów ds. młodzieży, którzy oceniają wykorzystanie sztucznej inteligencji w szkołach. Podkreśla ona napięcie między innowacją a nadzorem, pokazując, jak niewystarczające testowanie może podważyć zaufanie i wyniki nauczania. Podnosi również kluczowe kwestie dotyczące odpowiedzialności, uczciwości i dojrzałości narzędzi sztucznej inteligencji wykorzystywanych do oceny pracy uczniów. Młodzi ludzie są bezpośrednio dotknięci, gdy narzędzia algorytmiczne przypisują punkty, które kształtują informacje zwrotne, oceny, a nawet pewność siebie. Ten przypadek przypomina nam, że przejrzystość, ludzka ocena i zrozumienie kontekstu pozostają kluczowe, gdy sztuczna inteligencja jest wykorzystywana w kształtującej lub podsumowującej ocenie edukacyjnej.

Źródła

Studium przypadku 1:

- Axios:
<https://www.axios.com/2020/08/19/england-exams-algorithm-grading>
- Wired:
<https://www.wired.com/story/skewed-grading-algorithms-fuel-backlash-beyond-classroom>

Studium przypadku 2:

- Reuters – "Amazon scraps secret AI recruiting tool that showed bias against women" (2018):
<https://www.reuters.com/article/world/insight-amazon-scrap-secret-ai-recruiting-tool-that-showed-bias-against-women-idUSKCN1MK0AG/>
- LinkedIn – "Amazon's Biased Hiring Algorithm - A Wake-up Call for Tech" (2024)
<https://www.linkedin.com/pulse/amazons-biased-hiring-algorithm-2018-wake-up-call-tech-akshay-aryan-dfouc>
- The Independent – "Amazon scraps 'sexist AI' recruitment tool" (2018)
<https://www.independent.co.uk/tech/amazon-ai-sexist-recruitment-tool-algorithm-a8579161.html>
- Leoforce – "Lessons from Amazon's Sexist AI Recruiting Tool" (2024):
<https://leoforce.com/blog/amazon-ai-hiring-bias-case-study/>

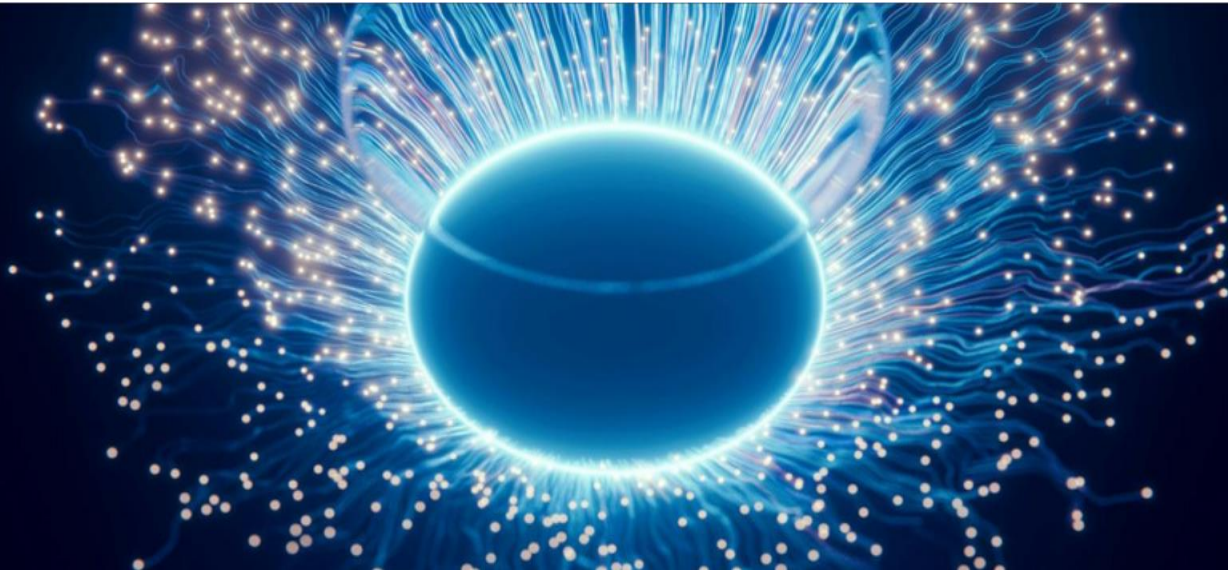
- Carnegie Mellon University – "Amazon Scraps Secret AI Recruiting Engine that Showed Biases Against Women" (2018)
<https://www.ml.cmu.edu/news/news-archive/2016-2020/2018/october/amazon-scrap-secret-artificial-intelligence-recruiting-engine-that-showed-biases-against-women.html>
- ACLU – "Why Amazon's Automated Hiring Tool Discriminated Against Women" (2018):
<https://www.aclu.org/news/womens-rights/why-amazons-automated-hiring-tool-discriminated-against>
- IMD – "Amazon's sexist hiring algorithm could still be better than a human" (2018):
<https://www.imd.org/research-knowledge/digital/articles/amazons-sexist-hiring-algorithm-could-still-be-better-than-a-human/>
- Redress Compliance – "Amazon AI Hiring Tool: A Case Study in Algorithmic Bias" (2025):
<https://redresscompliance.com/amazon-ai-hiring-tool-a-case-study-in-algorithmic-bias/>

Studium przypadku 3:

https://ceur-ws.org/Vol-3938/Paper_14.pdf
<https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2024.1221177/full>
Punoševac, Mina & Nikolić, Anđela. (2024). Transforming Assessment Practices – A Case Study of YouTestMe. CEUR Workshop Proceedings

Studium przypadku 4:

National Youth Assembly on AI Report (2023)
Government of Ireland, Department of Enterprise, Trade and Employment
<https://enterprise.gov.ie/en/publications/publication-files/national-youth-assembly-on-ai-report.pdf>



Źródła

Studium przypadku 5:

Milosevic et al. (2023). Effectiveness of Artificial Intelligence-Based Cyberbullying Interventions from a Youth Perspective
Dublin City University (DCU) Research, Ireland
<https://doras.dcu.ie/29411/1/milosevic-et-al-2023-effectiveness-of-artificial-intelligence-based-cyberbullying-interventions-from-youth-perspective.pdf>

Studium przypadku 6:

<https://arxiv.org/pdf/2409.07476>

https://www.researchgate.net/publication/383985608_Responsible_AI_for_Test_Equity_and_Quality_The_Duolingo_English_Test_as_a_Case_Study

Burstein, Jill & LaFlair, Geoffrey & Yancey, Kevin & von Davier, Alina & Dotan, Ravit. (2024).

Studium przypadku 7:

<https://globalwitness.org/en/campaigns/digital-threats/new-evidence-of-facebooks-sexist-algorithm/>

Opublikowano 12 czerwca 2023

Zaktualizowano 13 lutego 2025

Studium przypadku 8:

Frühwirth, M. & Lewerenz, T. (2024). A case study of an immersive learning unit for German as a second language lessons. Dostępne: <http://link.springer.com/article/10.1007/s44217-024-00106-w>

Studium przypadku 9:

Nöstlinger, N. (2024). Germany's far right is winning over the young. Politico.

<https://www.politico.eu/article/far-right-alternative-for-germany-afd-young-alternative-migration-remigration/>

Reuters. (2025). German spy agency brands far-right AfD as 'extremist'.

<https://www.reuters.com/world/europe/german-spy-agency-ranks-far-right-afd-extremist-2025-05-02/>

Politico. (2025). Germany's far-right AfD dissolves extremist youth branch to avert ban.

<https://www.politico.eu/article/afd-ends-youth-branch-germany/>

Le Monde diplomatique. (2025). Germany's AfD targets the youth vote.

<https://mondediplo.com/2025/01/09germany-afd>

Studium przypadku 10:

Röckl, Alexander & Bender, Maximilian (2024). "Wie gut ist die automatische Hausaufgabenbewertung mit KI?":

<https://arxiv.org/abs/2412.06651>



Wnioski

Jak już omawialiśmy w tym przewodniku, integracja sztucznej inteligencji w różnych obszarach społeczeństwa: edukacji, zatrudnieniu, mediach i innych, oferuje zarówno potencjał transformacyjny, jak i istotne wyzwania. Przedstawione studia przypadków pokazują różnorodne i kreatywne sposoby wykorzystania sztucznej inteligencji.

Łączy je wspólny mianownik – innowacja napędzana potrzebami i głosami młodych ludzi. Pokazują one, że przemyślane i etyczne wykorzystanie sztucznej inteligencji może pomóc młodym ludziom wziąć większą odpowiedzialność za swoją edukację, zatrudnienie i życie osobiste, a jednocześnie wspierać krytyczne myślenie o samej technologii.

Jednak te postępy podkreślają również znaczenie kompetencji cyfrowych, odpowiedzialnego korzystania z nich i równości dostępu. Edukatorzy, pracownicy młodzieżowi i decydenci – wszyscy mają do odegrania rolę w kształtowaniu wykorzystania sztucznej inteligencji w sposób, który priorytetowo traktuje relacje międzyludzkie, etyczne praktyki i uczenie się przez całe życie.

Ten przewodnik nie jest punktem końcowym, lecz zaproszeniem do dalszego eksplorowania, kwestionowania i współtworzenia przyszłości, w której młodzi ludzie nie będą jedynie użytkownikami sztucznej inteligencji, lecz aktywnymi twórcami jej roli w edukacji i społeczeństwie.

[Valuesai.eu](https://valuesai.eu)

www.valuesai.eu

przewodnik po studium przypadku 2025 VALUES jest
licencjonowany na podstawie licencji CC BY 4.0



śledź naszą podróż



Co-funded by
the European Union