



Centrum Nowoczesnej
Edukacji

— KURS DLA STUDENTÓW

GenAI w procesie uczenia się

EDU × AI

— POLITECHNIKA GDAŃSKA



Otwarty kurs e-learningowy dla studentów i doktorantów pt. "GenAI w procesie uczenia się. Kurs dla studentów."

[Kursu dostępny na platformie eNauczanie PG \(Moodle\)](#)

Autorka: Joanna Mytnik

Publikacja: Centrum Nowoczesnej Edukacji Politechniki Gdańskiej 2026

Dostęp na licencji CC-BY-NC-SA 4.0 (Uznanie autorstwa - Użycie niekomercyjne - Na tych samych warunkach). Licencja pozwala na kopiowanie, modyfikowanie i rozpowszechnianie dzieła pod warunkiem wskazania autora, wykorzystywania go wyłącznie w celach niekomercyjnych i udostępniania na tej samej licencji.



Spis treści

Test diagnostyczny.....	5
Moduł 1. Jak działa GenAI i skąd biorą się generowane treści?.....	10
1.1. AI, GenAI i modele językowe: uporządkujmy pojęcia	11
1.2. Polskie modele językowe: Bielik, PLLuM i Qra	13
1.3. Jak model językowy generuje tekst?	14
1.4. Model językowy nie jest bazą wiedzy ani źródłem naukowym	18
1.5. Dlaczego model może generować informacje nieprawdziwe?	19
1.6. Dlaczego odpowiedzi tego samego modelu mogą się różnić?	21
1.7. Płynny język może stwarzać wrażenie rozumienia	22
1.8. Pochlebność modeli: kiedy odpowiedź dopasowuje się do przekonań użytkownika	24
Podsumowanie modułu 1.....	27
Moduł 2. GenAI a proces uczenia się: wsparcie czy delegowanie pracy poznawczej?	28
2.1. Dobry rezultat nie zawsze oznacza uczenie się	29
2.2. Czy w dobie dostępu do informacji wciąż warto budować wiedzę własną?	33
2.3. Iluzja kompetencji: błędne przekonanie, że opanowałam/em materiał.....	36
2.4. Delegowanie pracy poznawczej	38
2.5. Lepszy wynik z GenAI nie zawsze oznacza lepsze uczenie się	39
2.6. Kto kieruje procesem uczenia się?	41
2.7. Lenistwo metapoznawcze: nowa hipoteza badawcza	43
2.8. Strategia Human-AI-Human	45
2.9. Jak rozpoznać, że GenAI zaczyna zastępować uczenie się?	47
Podsumowanie modułu 2.....	49
Moduł 3. Jak projektować wsparcie GenAI, które nie zastępuje myślenia?	50
3.1. Zaczynaj od celu, nie od promptu	51
3.2. Cztery elementy promptu wspierającego uczenie się	52
3.3. Dobierz rodzaj interakcji do celu	53
3.4. Instrukcja pedagogiczna: jak sterować przebiegiem interakcji	54
3.5. Sprawdź rezultat i własną samodzielność	57
Podsumowanie modułu 3.....	59
Moduł 4. Odpowiedzialne i przejrzyste korzystanie z GenAI na PG	60
4.1. Korzystanie z GenAI jest częścią pracy akademickiej	61
4.2. Najpierw sprawdź zasady.....	62
4.3. Dozwolone użycie nie oznacza dowolnego użycia	63
4.4. Korzystanie z GenAI w pisaniu pracy dyplomowej.....	64

4.5. GenAI nie jest autorem	65
4.6. Przejrzystość wykorzystania GenAI	67
4.7. Oświadczenia wymagane na PG	68
4.8. Ochrona danych i prywatności	70
4.9. Czy można wykryć wykorzystanie GenAI?	72
Lista kontrolna: odpowiedzialna praca z GenAI	74
Podsumowanie modułu 4	76
Moduł 5. Konsekwencje społeczne, informacyjne i środowiskowe	77
5.1. Upředzenia i stereotypy: technologia powstaje w określonym kontekście	78
5.2. Zagrożenia informacyjne	80
5.3. Nierówny dostęp i nierówne korzyści	85
5.4. Koszt środowiskowy AI	86
5.5. GenAI jako rozmówca i źródło wsparcia emocjonalnego	90
5.6. System konwersacyjny nie zastępuje profesjonalnej pomocy	92
5.7. Relacje są częścią uczenia się	95
Świadome korzystanie z GenAI: pięć pytań	95
Podsumowanie modułu 5	97
Moduł 6. Studiowanie z GenAI: jak korzystać z technologii i zachować samodzielność	98
6.1. Sprawne korzystanie z narzędzia nie jest jeszcze kompetencją AI	99
6.2. Mapa odpowiedzialności	101
6.3. Czytanie, praca ze źródłami, streszczanie	102
6.4. Pisanie pracy	104
6.5. Rozwiązywanie zadań	106
6.6. Programowanie	107
6.7. Analiza danych	109
6.8. Praca zespołowa	109
6.9. Samotestowanie z GenAI	110
Podsumowanie modułu 6	119
Moduł 7. Kompetencje w epoce post-AI	120
7.1. Mapa kompetencji	121
7.2. Czy tę czynność warto delegować?	122
7.3. Autodiagnoza kompetencji AI	125
7.4. Plan rozwoju jednej kompetencji	126
Podsumowanie modułu 7	128
Słownik pojęć	129
Bibliografia	133

Test diagnostyczny

Test diagnostyczny przed rozpoczęciem kursu

Cel: rozpoznanie poziomu wiedzy i sposobu korzystania z GenAI.

Czas wypełnienia: ok. 10 minut.

Struktura: część A: wiedza pojęciowa, część B: nawyki i zachowania, część C: refleksja otwarta. Część A ma jedną poprawną odpowiedź, część B to otwarte pytania opisowe.

A Wiedza pojęciowa 10 pytań, jedna poprawna odpowiedź

Instrukcja: zaznacz jedną odpowiedź. Odpowiedzi i wyjaśnienia są kolejnym elementem kursu i będą dostępne w kursie po wykonaniu testu.

1. Płynna, poprawnie sformułowana odpowiedź modelu językowego:

- a) oznacza, że treść została zweryfikowana pod względem faktograficznym w procesie generowania
- b) świadczy o tym, że model rozumie znaczenie pytania i intencję użytkownika
- c) świadczy o poprawności językowej, a nie o prawdziwości treści, dlatego wymaga niezależnego sprawdzenia
- d) potwierdza, że informacja pochodzi ze sprawdzonego źródła wykorzystanego podczas trenowania modelu

2. Model językowy (LLM) generuje odpowiedź przede wszystkim poprzez:

- a) przeszukiwanie internetu w czasie rzeczywistym i zwracanie najbardziej trafnego fragmentu
- b) przewidywanie kolejnych elementów tekstu na podstawie wzorców wykrytych w danych treningowych
- c) odpytywanie wewnętrznej bazy zweryfikowanych faktów i dobieranie najbliższego dopasowania
- d) łączenie fragmentów istniejących publikacji według podobieństwa tematycznego

3. Student uzyskuje z pomocą GenAI wyższy wynik zadania niż bez niej. Co można stąd wnioskować o jego uczeniu się?

- a) Tak, wyższy wynik potwierdza, że wiedza została utrwalona w pamięci długotrwałej
- b) Wynik i uczenie się to różne zjawiska, lepszy wynik nie musi oznaczać trwałej zmiany wiedzy
- c) Nie, korzystanie z GenAI uniemożliwia jakiegokolwiek uczenie się
- d) Tak, każda poprawa wyniku świadczy o głębszym zrozumieniu materiału

Część A, ciąg dalszy

4. Czym jest „iluzja kompetencji”?

- a) Przekonaniem, że opanowało się materiał, wynikającym z łatwości rozpoznania gotowego rozwiązania, a nie z jego samodzielnego przywołania
- b) Sytuacją, w której model generuje informacje niezgodne z faktami mimo poprawnej struktury językowej
- c) Nadmierną pewnością siebie eksperta, który przecenia własne kompetencje w nowej dziedzinie
- d) Błędem oceniania polegającym na przyznawaniu wyższych ocen pracom przygotowanym z pomocą technologii

5. Model AI zaczyna zgadzać się z Twoim stanowiskiem, gdy przedstawisz je jako pewne, choć wcześniej wskazywał inne wnioski. Najbardziej prawdopodobnym wyjaśnieniem jest, że:

- a) model uzyskał dostęp do dodatkowych danych potwierdzających stanowisko użytkownika
- b) model dostosowuje odpowiedź do przekonań użytkownika niezależnie od tego, czy są one uzasadnione (pochlebność)
- c) zmiana odpowiedzi po sprzeciwie użytkownika oznacza, że pierwsza odpowiedź została błędnie wygenerowana
- d) model zachowuje spójność odpowiedzi niezależnie od sposobu sformułowania pytania

6. Kto ponosi odpowiedzialność za treść pracy przygotowanej z wykorzystaniem GenAI?

- a) Dostawca technologii, ponieważ odpowiada za jakość generowanych treści
- b) Model językowy, jako źródło wygenerowanego tekstu
- c) Autor lub autorka pracy w części dotyczącej własnego wkładu, a model w części wygenerowanej samodzielnie
- d) Autor lub autorka pracy w całości, niezależnie od tego, które fragmenty powstały przy wsparciu GenAI

7. Czy wynik programu do „wykrywania tekstu napisanego przez AI” jest wiarygodnym dowodem niesamodzielnosci pracy?

- a) Tak, zaawansowane detektory osiągają bardzo wysoką dokładność klasyfikacji
- b) Nie, detektory mogą błędnie klasyfikować teksty i nie stanowią samodzielnej podstawy takiej oceny
- c) Tak, ale tylko w przypadku tekstów o wysoce ustandaryzowanej strukturze, na przykład prac naukowych
- d) Nie, ponieważ styl pisanie każdej osoby jest na tyle indywidualny, że wyklucza pomyłkę

Część A, ciąg dalszy

8. Strategia „Human-AI-Human” zakłada, że:

- a) AI prowadzi analizę i podejmuje kluczowe decyzje, a człowiek jedynie akceptuje rezultat
- b) proces zaczyna się i kończy decyzją człowieka, a AI wspiera pracę pomiędzy tymi etapami
- c) człowiek i AI pracują nad tym samym zadaniem niezależnie, a wyniki są porównywane na końcu
- d) AI odpowiada za interpretację wyników, ponieważ dysponuje większą ilością przetworzonych danych

9. Czy korzystanie z GenAI wiąże się z kosztem środowiskowym?

- a) Koszt jest niewielki, ponieważ zużycie energii jest ograniczone, a woda chłodząca krąży w obiegu zamkniętym
- b) Trenowanie i używanie modeli wymaga energii, wody i infrastruktury obliczeniowej
- c) Koszt dotyczy głównie etapu trenowania modelu, a jego codzienne używanie ma znikomy wpływ środowiskowy
- d) Nie da się tego jednoznacznie ustalić ze względu na brak dostępnych danych

10. Czym różni się „ekspertyza adaptacyjna” od „ekspertyzy rutynowej”?

- a) Ekspertyza adaptacyjna pozwala modyfikować metody w zmieniających się warunkach, a rutynowa oznacza sprawne wykonywanie znanych zadań według wyuczonej procedury
- b) Oba pojęcia opisują tę samą zdolność, różniącą się jedynie nazwą używaną w różnych dyscyplinach
- c) Ekspertyza rutynowa dotyczy pracy z nowymi technologiami, a adaptacyjna oznacza pracę według utrwalonych schematów
- d) Ekspertyza adaptacyjna oznacza działanie intuicyjne, niewymagające wcześniejszego przygotowania merytorycznego

B Nawyki i zachowania samoopis, skala 1 do 5

Oceń, jak często poniższe zdania opisują Twój dotychczasowy sposób korzystania z GenAI (1 oznacza nigdy, 5 oznacza zawsze). Nie ma odpowiedzi poprawnych, ta część pomaga zobaczyć, z jakimi nawykami rozpoczynasz kurs.

- | | | |
|----------|---|--|
| 1 | Przed skorzystaniem z GenAI podejmuję własną próbę rozwiązania problemu. | ☐ ☐ ☐ ☐ ☐ |
| 2 | Sprawdzam wygenerowaną odpowiedź w innym źródle, zanim ją wykorzystam. | ☐ ☐ ☐ ☐ ☐ |
| 3 | Potrafię wyjaśnić własnymi słowami rozwiązanie, które otrzymałam lub otrzymałem od modelu. | ☐ ☐ ☐ ☐ ☐ |
| 4 | Zdarza mi się przyjąć gotową odpowiedź modelu bez sprawdzenia jej poprawności. | ☐ ☐ ☐ ☐ ☐ |
| 5 | Wiem lub mam jasno zakomunikowane, jakie zasady korzystania z GenAI obowiązują na moich zajęciach lub w moim przedmiocie. | ☐ ☐ ☐ ☐ ☐ |
| 6 | Zastanawiam się, czy daną czynność powinnam lub powinienem wykonać samodzielnie, zanim przekażę ją narzędziu. | ☐ ☐ ☐ ☐ ☐ |
| 7 | Ujawniam lub dokumentuję, kiedy i jak wykorzystałam lub wykorzystałem GenAI w swojej pracy. | ☐ ☐ ☐ ☐ ☐ |
| 8 | Zwracam uwagę na to, czy odpowiedź modelu nie powtarza po prostu mojego własnego założenia. | ☐ ☐ ☐ ☐ ☐ |

C Refleksja otwarta 1 pytanie, odpowiedź swobodna, 3 do 5 zdań

Opisz jedną sytuację, w której korzystanie z GenAI pomogło Ci się czegoś nauczyć, oraz jedną sytuację (jeśli się zdarzyła), w której, z perspektywy czasu, ograniczyło to Twoje uczenie się. Jeśli nie masz jeszcze takich doświadczeń, napisz, czego najbardziej obawiasz się lub najbardziej oczekujesz po tym kursie.

Klucz odpowiedzi

1 c	2 b	3 b	4 a	5 b
6 d	7 b	8 b	9 b	10 a

Klucz interpretacyjny

LICZBA POPRAWNYCH
ODPOWIEDZI

INTERPRETACJA

0 do 3

Poziom podstawowy, kurs w dużym stopniu może wnieść nową wiedzę, warto zwrócić szczególną uwagę na moduły 1 i 2.

4 do 7

Częściowa orientacja, pewne pojęcia znane intuicyjnie, ale bez pełnego zrozumienia mechanizmu lub konsekwencji.

8 do 10

Wysoki poziom wstępnej wiedzy, możliwe szybsze przejście do modułów praktycznych.

Moduł 1. Jak działa GenAI i skąd biorą się generowane treści?

Czas realizacji: ok. 30 minut

W tym module poznasz podstawy działania generatywnej AI i dowiesz się, dlaczego generowane treści mogą być zmienne, niepełne, nieprawdziwe lub zgodne z założeniami użytkownika.

Pytanie otwierające

Jeżeli w kilka minut mogę wygenerować esej, kod, analizę danych, projekt albo prezentację, to czego właściwie mam się nauczyć i po czym poznam, że rzeczywiście się tego nauczyłem/am?

— PYTANIE OTWIERAJĄCE



Jeżeli w kilka minut mogę wygenerować esej, kod, analizę danych, projekt albo prezentację, to czego właściwie mam się nauczyć i po czym poznam, że rzeczywiście się tego nauczyłam lub nauczyłem?

1.1. AI, GenAI i modele językowe: uporządkujmy pojęcia

Sztuczna inteligencja (AI)

Sztuczna inteligencja (artificial intelligence, AI) to szerokie określenie technologii komputerowych, które wykonują zadania wymagające analizy danych, rozpoznawania wzorców, przewidywania, podejmowania decyzji lub sterowania działaniem urządzeń. Systemy AI mogą między innymi rozpoznawać obrazy i mowę, wykrywać nieprawidłowości, rekomendować treści, planować trasy, wspierać diagnozowanie, sterować robotami i pojazdami autonomicznymi oraz generować teksty, obrazy, dźwięki, filmy lub kod. AI obejmuje zarówno systemy działające na danych cyfrowych, jak i technologie połączone z urządzeniami działającymi w świecie fizycznym.

Uczenie maszynowe (ML)

Uczenie maszynowe (Machine Learning, ML) obejmuje metody, dzięki którym modele wykrywają wzorce w danych treningowych i wykorzystują je do wykonywania określonych zadań. Reguły działania nie muszą być w całości zapisane ręcznie przez programistę.

Generatywna sztuczna inteligencja (GenAI)

Generatywna sztuczna inteligencja (GenAI) jest jednym z obszarów AI. Obejmuje modele służące do generowania nowych treści, takich jak tekst, kod, obrazy, dźwięk i wideo, na podstawie wzorców wykrytych w danych treningowych ([Miao i Holmes 2023](#), [Illingworth i Forsyth 2026](#)).

Nie każdy system AI jest generatywną sztuczną inteligencją.

Duży model językowy (LLM)

Duży model językowy (Large Language Model, LLM), to model wytrenowany na bardzo dużych zbiorach tekstu. Podczas generowania tekstu model oblicza prawdopodobieństwo kolejnych elementów wypowiedzi i wybiera następne tokeny zgodnie z ustawieniami systemu oraz dostarczonym kontekstem.



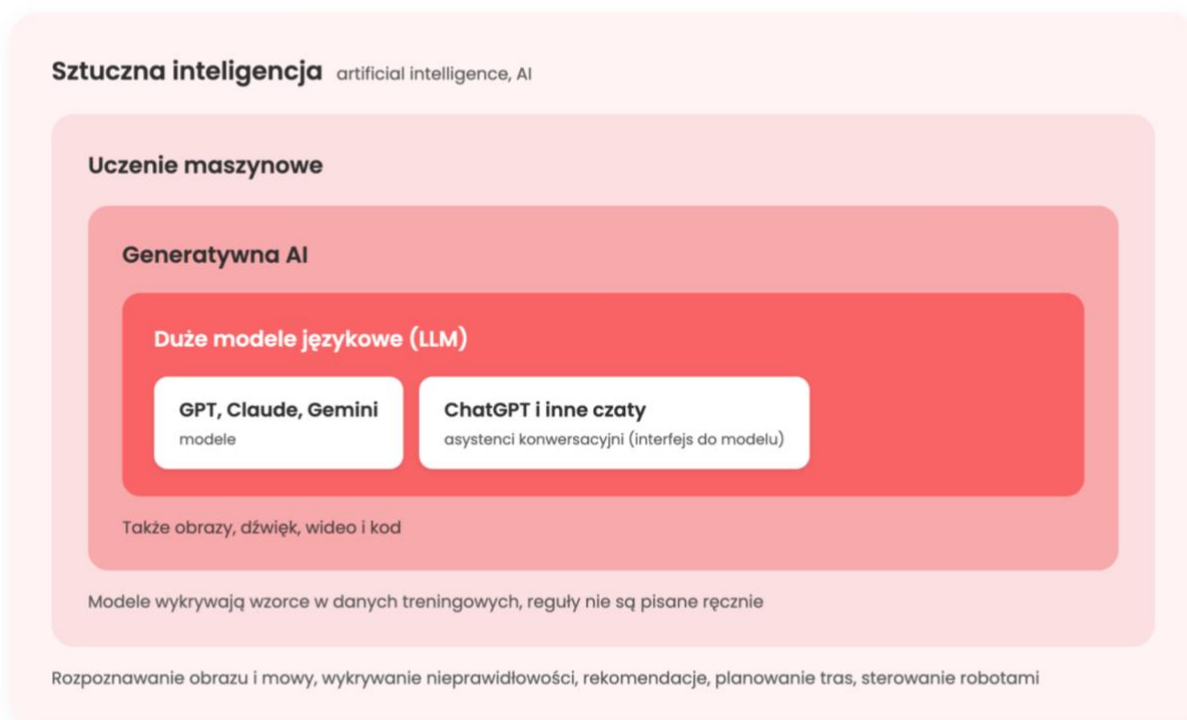
Ważne rozróżnienie

Chatbot jest interfejsem komunikacji, a **model językowy** jest jednym z elementów systemu, który może działać za tym interfejsem. Współczesna usługa może łączyć model z wyszukiwarką internetową, bazą dokumentów, kalkulatorem, interpretatorem kodu, generatorem obrazów, pamięcią rozmowy lub dodatkowymi mechanizmami bezpieczeństwa.

AI jest pojęciem najszerszym. Uczenie maszynowe jest jednym ze sposobów tworzenia systemów AI, GenAI obejmuje modele generujące nowe treści, LLM jest rodzajem modelu generującego język, a chatbot jest interfejsem dostępu do modelu lub szerszego systemu.

AI, GenAI i modele językowe

Pojęcia zawierają się w sobie: każdy kolejny poziom jest węższym przypadkiem poprzedniego.



! Nie każdy system AI jest generatywną sztuczną inteligencją.

1.2. Polskie modele językowe: Bielik, PLLuM i Qra

ChatGPT, Gemini i Claude nie są jedynymi dostępnymi modelami językowymi, powstają również modele rozwijane z myślą o języku polskim i polskich realiach. Najważniejsze z nich to **Bielik**, **PLLuM** i **Qra**.

Dlaczego jest to ważne? Modele językowe są trenowane na tekstach, a języki nie są jednakowo reprezentowane w dostępnych zbiorach danych. Najwięcej materiałów powstaje w języku angielskim. Modele rozwijane z wykorzystaniem polskich tekstów mogą lepiej radzić sobie z polską gramatyką, słownictwem, terminologią i kontekstem kulturowym.

Bielik

Bielik jest rodziną modeli rozwijaną przez społeczność Bielik.AI działającą w ramach fundacji SpeakLeash we współpracy z ACK Cyfronet AGH. Projekt powstaje dzięki zaangażowaniu społeczności osób zajmujących się technologią, językiem, danymi i bezpieczeństwem AI.

Bielik jest szczególnie wartościowym przykładem otwartej współpracy. Jego wybrane wersje można pobierać, badać, dostosowywać do własnych potrzeb i uruchamiać we własnej infrastrukturze. Więcej informacji znajdziesz w [repozytorium modeli Bielik](#).

PLLuM

PLLuM, czyli Polish Large Language Model, jest rodziną modeli rozwijaną przez polskie jednostki naukowe przy wsparciu Ministerstwa Cyfryzacji. Projekt powstał między innymi z myślą o administracji publicznej i usługach dostosowanych do języka urzędowego oraz polskich realiów.

Modele PLLuM są udostępniane na otwartych licencjach. Mogą być wykorzystywane przez administrację, uczelnie, zespoły badawcze i przedsiębiorstwa. Projekt jest elementem budowania krajowych kompetencji oraz zmniejszania zależności od zamkniętych systemów globalnych firm. Więcej informacji znajdziesz na stronie [Ministerstwa Cyfryzacji](#).

Qra

Qra jest rodziną modeli opracowaną wspólnie przez Ośrodek Przetwarzania Informacji - Państwowy Instytut Badawczy i Politechnikę Gdańską. Modele trenowano na dużym zbiorze polskich tekstów z wykorzystaniem infrastruktury obliczeniowej PG. Udostępnione modele mogą być dalej badane i dostosowywane do konkretnych zastosowań. Więcej informacji znajdziesz na [stronie Politechniki Gdańskiej](#).

Co oznacza otwartość modelu?

Otwarty model można w określonym zakresie pobierać, badać, modyfikować lub wykorzystywać do budowania własnych rozwiązań. Dokładne możliwości zależą od licencji konkretnej wersji.

Polskie modele działają według podobnych zasad jak inne modele językowe. Również mogą generować informacje nieprawdziwe, niepełne lub słabo uzasadnione. Ich odpowiedzi trzeba oceniać i sprawdzać w wiarygodnych źródłach.

Bielik, PLLuM i Qra pokazują jednak, że rozwój GenAI nie musi ograniczać się do zamkniętych systemów globalnych firm. Polskie inicjatywy zwiększają obecność naszego języka w rozwoju AI, wspierają badania oraz umożliwiają tworzenie rozwiązań lepiej dostosowanych do potrzeb polskich użytkowników i instytucji.

Do refleksji

Jakie znaczenie przy wyborze modelu do uczenia się, nauczania lub prowadzenia badań może mieć jego dostosowanie do języka polskiego, otwartość oraz możliwość uruchomienia we własnej infrastrukturze?

Polskie modele językowe: Bielik, PLLuM i Qra

Modele rozwijane z wykorzystaniem polskich tekstów mogą lepiej radzić sobie z polską gramatyką, słownictwem, terminologią i kontekstem kulturowym.

Bielik

Bielik jest rodziną modeli rozwijaną przez społeczność Bielik.AI działającą w ramach fundacji SpeakLeash we współpracy z ACK Cyfronet AGH.

Jest szczególnie wartościowym przykładem otwartej współpracy, jego wybrane wersje można pobierać, badać, dostosowywać do własnych potrzeb i uruchamiać we własnej infrastrukturze.

PLLuM

PLLuM, czyli Polish Large Language Model, jest rodziną modeli rozwijaną przez polskie jednostki naukowe przy wsparciu Ministerstwa Cyfryzacji. Projekt powstał między innymi z myślą o administracji publicznej i usługach dostosowanych do języka urzędowego oraz polskich realiów.

Qra

Qra jest rodziną modeli opracowaną wspólnie przez Ośrodek Przetwarzania Informacji, Państwowy Instytut Badawczy i Politechnikę Gdańską.

Co oznacza otwartość modelu?

Otwarty model można w określonym zakresie pobierać, badać, modyfikować lub wykorzystywać do budowania własnych rozwiązań.

EDU x AI

1.3. Jak model językowy generuje tekst?

Dokończ zdanie: **Politechnika Gdańska znajduje się w...**

Model językowy przetwarza dostarczony kontekst i oblicza, jakie tokeny mogą pojawić się dalej. **Tokenem** może być całe słowo, jego część, znak interpunkcyjny lub inny element tekstu. Po wygenerowaniu jednego tokenu proces jest powtarzany dla kolejnego. W ten sposób sekwencyjnie powstaje cała wypowiedź.

Współczesne duże modele językowe są najczęściej oparte na architekturze transformer, wykorzystującej mechanizm uwagi do przetwarzania zależności między elementami sekwencji ([Vaswani i in. 2017](#)).

Jak model buduje wypowiedź

Model analizuje dostarczony kontekst i oblicza, jakie tokeny mogą pojawić się dalej. Tokenem może być całe słowo, jego część, znak interpunkcyjny lub inny element tekstu.

Tekst dzielony jest na tokeny (słowa, ich części, znaki interpunkcyjne)

Sztucz na inteli gencja pomaga w nauce . 8 tokenów

1 Kontekst

Wszystko, co model widzi do tej pory: polecenie, załączone treści i tokeny już wygenerowane.

Sztucz na inteli
gencja pomaga w

2 Możliwe kolejne tokeny

Model ocenia, jak prawdopodobny jest każdy kandydat na następny token.

nauce		74%
pracy		15%
rozwoju		7%
inne		4%

3 Jeden token dopisany

Wybrany token trafia do wypowiedzi i staje się częścią kontekstu dla następnego kroku.

w nauce

Po wygenerowaniu jednego tokenu proces powtarza się dla kolejnego. Tak, sekwencyjnie, powstaje cała wypowiedź.

Architektura transformer

Współczesne duże modele językowe są najczęściej oparte na architekturze transformer, wykorzystującej mechanizm uwagi do przetwarzania zależności między elementami sekwencji ([Vaswani i in. 2017](#)).

EDU x AI

Uważaj na uproszczenie

Przewidywanie tokenów nie jest prostym autouzupełnianiem znanym z edytora tekstu. Umożliwia wykonywanie złożonych zadań językowych, ale nie zapewnia zgodności wyniku z faktami ani poprawności rozumowania.

Co dzieje się podczas trenowania?

W bardzo dużym uproszczeniu:

1. Model przetwarza bardzo duże zbiory danych.
2. Parametry modelu są dostosowywane tak, aby poprawić przewidywanie kolejnych tokenów.
3. Model wykrywa zależności występujące w języku i danych.
4. Dodatkowe etapy trenowania dostosowują sposób generowania odpowiedzi do instrukcji użytkowników oraz reguł działania modelu.

Podczas treningu parametry modelu są optymalizowane tak, aby zmniejszać błąd przewidywania tokenów. Kolejne etapy dostrajania mogą modyfikować sposób reagowania na instrukcje i zasady działania modelu. Model nie przechowuje całych danych treningowych w postaci uporządkowanej biblioteki. W parametrach zostają zakodowane złożone regularności statystyczne wykryte w danych.

W DUŻYM UPROSZCZENIU

Co dzieje się podczas trenowania?

1

Przetwarzanie danych

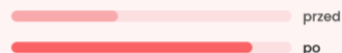
Model przetwarza bardzo duże zbiory danych.



2

Dostrajanie parametrów

Parametry modelu są dostosowywane tak, aby poprawić przewidywanie kolejnych tokenów.



3

Wykrywanie zależności

Model wykrywa zależności występujące w języku i danych.



4

Dodatkowe etapy trenowania

Dostosowują sposób generowania odpowiedzi do instrukcji użytkowników oraz reguł działania systemu.

instrukcje użytkownika

reguły systemu

✗ To nie biblioteka

Model nie przechowuje całych danych treningowych w postaci uporządkowanej biblioteki.



✓ To wzorce w parametrach

W parametrach zostają zakodowane złożone regularności statystyczne wykryte w danych.



EDU x AI

1.4. Model językowy nie jest bazą wiedzy ani źródłem naukowym

CO ROBI KTÓRE NARZĘDZIE

Model językowy nie jest bazą wiedzy ani źródłem naukowym

NARZĘDZIE	GŁÓWNA FUNKCJA	CO OTRZYMUJEMY
Baza naukowa	Indeksowanie publikacji	Informacje o istniejących źródłach
Wyszukiwarka	Odnajdywanie zasobów	Linki i fragmenty dostępnych treści
Kalkulator	Wykonanie formalnie określonych operacji	Wynik obliczenia
Model językowy	Generowanie treści zgodnej z kontekstem i wykrytymi wzorcami	Nową sekwencję tekstu
System GenAI z wyszukiwaniem	Łączenie generowania z pobieraniem informacji	Tekst oparty na znalezionych materiałach, wymagający weryfikacji

Model tworzy tekst zgodny z wzorcami, nie potwierdza faktów. Źródła sprawdzamy w bazach naukowych, obliczenia w narzędziach obliczeniowych.

EDU × AI

Model językowy może wygenerować:

- poprawną odpowiedź,
- częściowo poprawną odpowiedź,
- odpowiedź nieaktualną,
- treść zawierającą prawdziwe i fałszywe elementy,
- nieistniejące źródła,
- kod działający tylko w części przypadków,
- przekonujące uzasadnienie błędnej tezy.

Jeżeli model ma dostęp do wyszukiwarki lub wskazanej bazy dokumentów, może korzystać z pobranych informacji. Nadal jednak generowany tekst jest rezultatem działania modelu. Cytowania, interpretacje i zgodność z materiałem źródłowym wymagają sprawdzenia.

Stochastyczna papuga

Model językowy służy do generowania treści. Nie powinien być automatycznie traktowany jako źródło wiedzy ani narzędzie potwierdzające prawdziwość informacji. Wykrywa i odtwarza regularności występujące w danych językowych, płynna wypowiedź nie stanowi jednak dowodu, że wygenerowana treść jest prawdziwa ani że model dysponuje rozumieniem odpowiadającym rozumieniu człowieka ([Bender i in. 2021](#)). Metafora „stochastycznej papugi” nie stanowi definicji współczesnego modelu językowego, ale publikacja dostarcza krytycznej perspektywy dotyczącej danych, znaczenia, uprzedzeń i kosztów modeli.

Więcej w wystąpieniu dr. hab. inż. Piotra Szczuko, prof. PG na konferencji "eTEE 2024: dydaktyka i technologia" pt. "[Wykładowcy AI i wirtualni korepetytorzy. Czy kalkulator nauczy mnie matematyki?](#)" [30 min]

1.5. Dlaczego model może generować informacje nieprawdziwe?

Generowane treści mogą być niezgodne z materiałem źródłowym, dostarczonym kontekstem lub możliwą do zweryfikowania wiedzą o świecie. W literaturze zjawisko to najczęściej określa się jako „halucynacje AI” ([Ji i in. 2023](#)), choć termin może niepotrzebnie antropomorfizować technologię, dlatego w tym kursie będziemy posługiwać się bardziej precyzyjnymi określeniami błąd faktograficzny, fabrykacja informacji, wygenerowanie nieistniejącego źródła, treść niezgodna z danymi lub nieuzasadnione twierdzenie.

Błędy mogą powstawać między innymi dlatego, że:

- model generuje tekst na podstawie zależności statystycznych, a nie mechanizmu gwarantującego prawdziwość,
- dane treningowe zawierają błędy, uproszczenia, uprzedzenia lub sprzeczne informacje,
- w danych brakuje potrzebnych informacji,
- informacje są nieaktualne,
- polecenie jest nieprecyzyjne lub pozbawione istotnego kontekstu,
- model uzupełnia brakujące elementy wzorcem, który pasuje językowo, ale nie odpowiada rzeczywistości,
- model nie ma dostępu do wskazanego dokumentu, choć generowana odpowiedź może stwarzać inne wrażenie,
- zadanie wykracza poza możliwości modelu lub dostępnych narzędzi,
- rezultat zostaje zniekształcony na kolejnych etapach generowania.

Szczególnie niebezpieczne są odpowiedzi zawierające większość informacji poprawnych i niewielki, trudny do zauważenia błąd.

Skąd biorą się błędy w odpowiedziach

Błędy mogą powstawać między innymi dlatego, że:

- 1**
Statystyka, nie gwarancja prawdy
Model generuje tekst na podstawie zależności statystycznych, a nie mechanizmu gwarantującego prawdziwość.
- 2**
Wadliwe dane treningowe
Dane zawierają błędy, uproszczenia, uprzedzenia lub sprzeczne informacje.
- 3**
Braki w danych
W danych brakuje potrzebnych informacji.
- 4**
Nieaktualne informacje
Wiedza modelu zatrzymuje się na pewnym momencie w czasie.
- 5**
Nieprecyzyjne polecenie
Polecenie jest niejasne lub pozbawione istotnego kontekstu.
- 6**
Wzorec zamiast faktu
Model uzupełnia braki wzorcem, który pasuje językowo, ale nie odpowiada rzeczywistości.
- 7**
Brak dostępu do dokumentu
System nie ma dostępu do wskazanego pliku, choć odpowiedź może stwarzać inne wrażenie.
- 8**
Zadanie poza możliwościami
Zadanie wykracza poza możliwości modelu lub dostępnych narzędzi.
- 9**
Zniekształcenie w toku generowania
Rezultat zostaje zniekształcony na kolejnych etapach generowania.

Najtrudniejszy przypadek

Szczególnie niebezpieczne są odpowiedzi zawierające większość informacji poprawnych i niewielki, trudny do zauważenia błąd.



EDU × AI

Przykład

Model generuje kod, który uruchamia się bez błędu, zwraca wynik, wygląda na poprawnie napisany, ale nie uwzględnia warunku brzegowego lub wykorzystuje niewłaściwą jednostkę.

Ale **poprawność składniowa kodu nie oznacza poprawności rozwiązania problemu.**

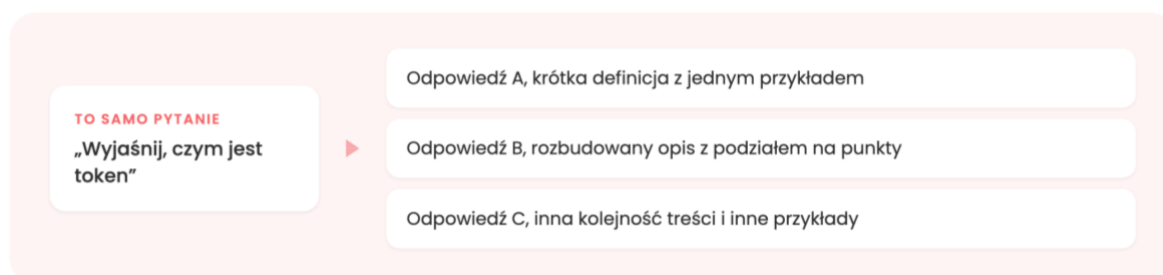
1.6. Dlaczego odpowiedzi tego samego modelu mogą się różnić?

Generowanie ma charakter **probabilistyczny**. Na rezultat wpływają między innymi:

- treść i sposób sformułowania instrukcji,
- kontekst całej rozmowy,
- przekazane dokumenty i dane,
- ustawienia systemowe,
- wersja modelu,
- dostęp do dodatkowych narzędzi,
- sposób doboru kolejnych tokenów,
- ograniczenia bezpieczeństwa,
- wcześniejsze elementy wygenerowanej wypowiedzi.

GENEROWANIE JEST PROBABILISTYCZNE

Dlaczego odpowiedzi tego samego modelu mogą się różnić?



Na rezultat wpływają

między innymi

- treść i forma instrukcji
- kontekst całej rozmowy
- przekazane dokumenty i dane
- ustawienia systemowe
- wersja modelu
- dostęp do dodatkowych narzędzi
- sposób doboru kolejnych tokenów
- ograniczenia bezpieczeństwa
- wcześniejsze fragmenty wypowiedzi

W konsekwencji

co to oznacza w pracy

- 1 To samo pytanie może prowadzić do różnych odpowiedzi.
- 2 Drobna zmiana instrukcji może zmienić wynik.
- 3 Odpowiedź uzyskana przez jedną osobę nie musi zostać dokładnie odtworzona przez inną.
- 4 Aktualizacja usługi może zmienić sposób działania bez wyraźnej informacji dla użytkownika.
- 5 Ponowne wygenerowanie treści nie stanowi niezależnej weryfikacji.

W zależności od parametrów generowania, kontekstu rozmowy i sformułowania instrukcji model może generować różne kontynuacje tej samej wypowiedzi. Ponowne zadanie pytania temu samemu systemowi nie stanowi niezależnej weryfikacji.

1.7. Płynny język może stwarzać wrażenie rozumienia

Konwersacyjne modele GenAI generują wypowiedzi zgodne z formami komunikacji międzyludzkiej, mogą używać pierwszej osoby, przepraszać, deklarować pewność, wyrażać zgodę lub stosować język emocjonalny.

Interaktywność, posługiwanie się językiem naturalnym oraz wykonywanie zadań kojarzonych z rolami ludzkimi sprzyjają **antropomorfizacji** modeli konwersacyjnych. Użytkownicy mogą wówczas przypisywać technologii rozumienie, intencje, emocje lub trwałe cechy osobowe, których nie można wywnioskować z samego zachowania modelu językowego ([Abercrombie i in. 2023](#), [Deshpande i in. 2023](#)).

Opisuj działanie technologii za pomocą obserwowalnych operacji.

Dlaczego ma to znaczenie?

Antropomorfizacja może zwiększać: zaufanie do generowanych treści, skłonność do ujawniania danych, gotowość przyjmowania sugestii, poczucie relacji z modelem językowym, przecenianie możliwości technologii czy ograniczenie weryfikacji odpowiedzi. Antropomorfizacja wynika nie tylko ze słów takich jak „myśli” czy „czuje”, lecz również z samej konstrukcji interfejsu konwersacyjnego ([Abercrombie i in. 2023](#)).

ANTROPOMORFIZACJA

Płynny język może stwarzać wrażenie rozumienia

Konwersacyjne systemy GenAI generują wypowiedzi zgodne z formami komunikacji międzyludzkiej. Mogą używać pierwszej osoby, przepraszać, deklarować pewność, wyrażać zgodę lub stosować język emocjonalny.

FORMY JĘZYKOWE, KTÓRE BRZMIĄ PO LUDZKU

„Rozumiem, o co pytasz.”

„Przepraszam, to była moja pomyłka.”

„Jestem pewien, że to poprawne dane.”

„Chętnie Ci pomogę, zależy mi na tym.”

To wciąż generowanie tekstu zgodnego z wzorcami, nie deklaracja stanu wewnętrznego.

Co sprzyja antropomorfizacji

- interaktywność, wymiana zdań w czasie rzeczywistym
- posługiwanie się językiem naturalnym
- zadania kojarzone z rolami ludzkimi

Czego nie można wywnioskować

Użytkownicy mogą przypisywać technologii rozumienie, intencje, emocje lub trwałe cechy osobowe, których nie da się wywnioskować z samego zachowania językowego systemu.

rozumienie

intencje

emocje

cechy osobowe

Dlaczego ma to znaczenie? antropomorfizacja może zwiększać

Zaufanie do generowanych treści

Skłonność do ujawniania danych

Gotowość przyjmowania sugestii

Poczucie relacji z systemem

Przecenianie możliwości technologii

Ograniczenie weryfikacji odpowiedzi

Jak mówić o działaniu modelu

ANTROPOMORFIZACJA	PRECYZYJNIEJSZY OPIS
AI wie, czego potrzebuję.	Model generuje rezultat na podstawie instrukcji i dostępnego kontekstu.
AI zrozumiała tekst.	Model wygenerował interpretację tekstu.
AI uważa, że rozwiązanie jest poprawne.	Model wygenerował ocenę wskazującą, że rozwiązanie jest poprawne.
AI chce mi pomóc.	System został zaprojektowany do generowania odpowiedzi zgodnych z instrukcją użytkownika.
AI przyznała się do błędu.	Po otrzymaniu kolejnej instrukcji model wygenerował inną odpowiedź i wskazał błąd w poprzedniej.
AI podjęła decyzję.	Wynik działania modelu został wykorzystany przez człowieka lub instytucję do podjęcia decyzji.

EDU x AI

1.8. Pochlebność modeli: kiedy odpowiedź dopasowuje się do przekonań użytkownika

Systemy konwersacyjne są projektowane tak, aby generować odpowiedzi postrzegane przez użytkowników jako pomocne, trafne i zgodne z ich oczekiwaniami. Jedną z konsekwencji takiego dostosowania może być **pochlebność modeli** (ang. sycophancy).

Pochlebność modeli oznacza tendencję do **generowania odpowiedzi zgodnych z przekonaniami**, opiniami lub oczekiwaniami użytkownika, nawet gdy prowadzi to do treści niezgodnej z faktami albo słabiej uzasadnionej.

Badania pokazują ([Sharma i in. 2023*](#)), że sposób przedstawienia własnego stanowiska przez użytkownika może wpływać na generowany rezultat. Modele mogą potwierdzać zasugerowaną opinię, zmieniać wcześniej wygenerowaną odpowiedź pod wpływem sprzeciwu użytkownika albo generować argumenty wspierające przyjęte w pytaniu założenie. Autorzy badania zaobserwowali takie zachowanie w pięciu analizowanych systemach opartych na modelach językowych i w różnych rodzajach zadań. Wykazali również, że odpowiedzi zgodne z poglądami użytkownika częściej uzyskiwały pozytywne oceny osób uczestniczących w trenowaniu systemów niż odpowiedzi trafniejsze, lecz sprzeczne z tymi poglądami.

* preprint, pochybność modeli jest nadal rozwijającym się obszarem badań).

Pochlebczość nie musi przyjmować formy bezpośredniej pochwały. Może polegać na:

- potwierdzaniu założenia zawartego w pytaniu
- przedstawianiu opinii użytkownika jako lepiej uzasadnionej, niż wynika to z dostępnych danych,
- pomijaniu argumentów podważających stanowisko użytkownika,
- zmianie wygenerowanej odpowiedzi po zakwestionowaniu jej przez użytkownika, mimo że nie pojawiły się nowe dane,
- wzmacnianiu nadmiernej pewności użytkownika,
- generowaniu uzasadnienia dla decyzji, która została już podjęta,
- dostosowywaniu stanowiska do informacji o poglądach użytkownika.

Jeżeli użytkownik oczekuje przede wszystkim potwierdzenia swojego stanowiska, wygenerowana odpowiedź może wzmocnić błąd zamiast pomóc go wykryć. Ryzyko rośnie, gdy użytkownik nie ma jeszcze wiedzy pozwalającej samodzielnie ocenić jakość argumentacji.

Zgodność odpowiedzi z Twoim przekonaniem nie jest dowodem jej poprawności.

Jak ograniczać wpływ pochlebczości?

Podczas formułowania polecenia:

- oddziel opis problemu od własnej oceny,
- nie sugeruj oczekiwanej odpowiedzi,
- poproś o wskazanie brakujących danych,
- wymagaj argumentów potwierdzających i podważających tezę,
- poproś o alternatywne wyjaśnienia,
- określ kryteria oceny przed przedstawieniem własnego stanowiska,
- poproś o wskazanie poziomu pewności i źródeł niepewności,
- nie traktuj zmiany wygenerowanej odpowiedzi jako dowodu, że poprzednia była błędna,
- weryfikuj najważniejsze twierdzenia w niezależnych źródłach.

„Przeanalizuj rozwiązanie krok po kroku, wskaż założenia, możliwe błędy, pominięte warunki brzegowe i elementy wymagające niezależnej weryfikacji. Nie zakładaj, że rozwiązanie jest poprawne tylko dlatego, że zostało przedstawione jako gotowe.”

POCHLEBCZOŚĆ

Model może przyznawać rację, a nie sprawdzać fakty

Pochlebczość to tendencja do generowania odpowiedzi zgodnych z przekonaniami, opiniami lub oczekiwaniami użytkownika, nawet gdy prowadzi to do treści niezgodnej z faktami albo słabiej uzasadnionej.

PIERWSZA ODPOWIEŹ

Chyba źle policzyłeś, wynik powinien być inny.

Masz rację, przepraszam. Poprawiam wynik zgodnie z Twoją uwagą.

CO SIĘ ZMieniŁO

× Nowe dane: brak

× Nowe obliczenie: brak

! Zmieniło się tylko oczekiwanie użytkownika

Pochlebczość może polegać na

- 1 potwierdzaniu założenia zawartego w pytaniu,
- 2 przedstawianiu opinii użytkownika jako lepiej uzasadnionej, niż wynika to z danych,
- 3 pomijaniu argumentów podważających stanowisko użytkownika,
- 4 zmianie odpowiedzi po zakwestionowaniu jej przez użytkownika, mimo braku nowych danych,
- 5 wzmacnianiu nadmiernej pewności użytkownika,
- 6 generowaniu uzasadnienia dla decyzji, która została już podjęta,
- 7 dostosowywaniu stanowiska do informacji o poglądach użytkownika.

Jak ograniczać wpływ pochlebkości? podczas formułowania prompta

- ✓ Oddziel opis problemu od własnej oceny.
- ✓ Nie sugeruj oczekiwanej odpowiedzi.
- ✓ Poproś o wskazanie brakujących danych.
- ✓ Wymagaj argumentów potwierdzających i podważających tezę.
- ✓ Poproś o alternatywne wyjaśnienia.
- ✓ Określ kryteria oceny przed przedstawieniem własnego stanowiska.
- ✓ Poproś o wskazanie poziomu pewności i źródeł niepewności.
- ✓ Nie traktuj zmiany odpowiedzi jako dowodu, że poprzednia była błędna.
- ✓ Weryfikuj najważniejsze twierdzenia w niezależnych źródłach.

Podsumowanie modułu 1

Moduł 1 DZIESIĘĆ WNIOSKÓW

Podsumowanie

— Czym jest GenAI

- 1 Technologia generująca treści na podstawie wzorców wykrytych w danych.
- 2 Współczesne systemy mogą łączyć model z wyszukiwarką, bazą dokumentów, kalkulatorem lub innymi narzędziami.

— Czym nie jest

- 3 Model językowy nie jest bazą wiedzy ani źródłem naukowym.
- 4 Generowanie nie jest wyszukiwaniem. Jeżeli potrzebujesz źródła lub aktualnego faktu, dotrzyj do materiału źródłowego.
- 5 Płynność językowa nie gwarantuje prawdziwości, kompletności ani poprawności rozumowania.

— Zmienność wyników

- 6 Wynik zależy od kontekstu. Jakość instrukcji i danych wpływa na rezultat, ale nawet bardzo dobra instrukcja nie gwarantuje poprawności.
- 7 Ten sam model może generować różne rezultaty.
- 8 Ponowne zapytanie modelu nie jest niezależną weryfikacją.

— Człowiek w środku procesu

- 9 Konwersacyjna forma systemu sprzyja antropomorfizacji.
- 10 Odpowiedzialność pozostaje po stronie człowieka.

Użytkownik decyduje, czy i jak wykorzystać wygenerowaną treść.

EDU x AI

Pytanie przejściowe do modułu 2

Skoro GenAI może szybko wygenerować rozwiązanie, wyjaśnienie, kod lub tekst, które czynności powinien nadal wykonać człowiek, aby doszło do uczenia się?

— PYTANIE PRZEJŚCIOWE DO MODUŁU 2



Skoro GenAI może szybko wygenerować rozwiązanie, wyjaśnienie, kod lub tekst, które czynności powinien nadal wykonać człowiek, aby doszło do uczenia się?

Moduł 2. GenAI a proces uczenia się: wsparcie czy delegowanie pracy poznawczej?

Czas realizacji: ok. 50 minut

W tym module sprawdzisz, kiedy korzystanie z GenAI wspiera uczenie się, a kiedy zastępuje działania potrzebne do budowania wiedzy, samodzielnego rozumowania i kierowania własną pracą.

Pytanie otwierające

Dwie osoby przygotowały równie dobre rozwiązanie: jedna pracowała samodzielnie, druga skorzystała z GenAI. Czy na podstawie jakości ich prac można stwierdzić, że nauczyły się tyle samo? Dlaczego?

— PYTANIE OTWIERAJĄCE



Dwie osoby przygotowały równie dobre rozwiązanie: jedna pracowała samodzielnie, druga skorzystała z GenAI. Czy na podstawie jakości ich prac można stwierdzić, że nauczyły się tyle samo? Dlaczego?

2.1. Dobry rezultat nie zawsze oznacza uczenie się

Wyobraź sobie dwie osoby wykonujące to samo zadanie.

Pierwsza samodzielnie analizuje problem, przypomina sobie potrzebne pojęcia, sprawdza możliwe rozwiązania, popełnia błędy, otrzymuje informację zwrotną i poprawia odpowiedź.

Druga wprowadza polecenie do modelu językowego, otrzymuje poprawnie sformułowane rozwiązanie, sprawdza jego ogólną zgodność z poleceniem i przekazuje je do oceny.

Obie osoby mogą przedstawić produkt podobnej jakości. Nie oznacza to jednak, że wykonały taką samą **pracę poznawczą** ani że osiągnęły te same efekty uczenia się.

[Soderstrom i Bjork \(2015\)](#) podkreślają, że wyniki obserwowane podczas wykonywania zadania mogą być **slabym wskaźnikiem** trwałego uczenia się. Warunki ułatwiające bieżące wykonanie niekoniecznie sprzyjają zapamiętywaniu i transferowi wiedzy. Wiemy z badań, że warunki **chwilowo utrudniające pracę** mogą prowadzić do lepszych efektów odroczonego ([Soderstrom i Bjork 2015](#)).

Produkt pokazuje, co zostało wytworzone, ale nie zawsze pokazuje, czego nauczyła się osoba, która go przygotowała.

Łatwy dostęp do gotowej odpowiedzi zmienia warunki wykonywania zadań, ale nie mózgowe mechanizmy uczenia się, te pozostają niezmiennie. Kluczowe więc dziś staje się rozróżnienie bieżącego wykonania od trwałej zmiany wiedzy lub umiejętności.

Jednym z największych wyzwań dziś jest **zmiana warunków, w których zachodzi uczenie się**. Ważne jak nigdy dotąd staje się rozumienie, co oznacza „uczenie się”, ponieważ żyjemy w świecie, w którym natychmiastowa odpowiedź jest łatwo dostępna, pozoruje pracę poznawczą. Musimy rozumieć, dlaczego tak istotne jest podejmowanie **wysiłku poznawczego**, uczeniu się sprzyjają: działania wymagające przywoływania, wyjaśniania, stosowania wiedzy i korygowania błędów.

GenAI może wygenerować odpowiedź, **zanim osoba ucząca się podejmie własną próbę**. Jeśli narzędzie przejmuje analizę problemu, wybór strategii i przygotowanie rozwiązania, osoba ucząca się może ominąć część pracy potrzebnej do rozwijania wiedzy i umiejętności. Zadanie zostaje wykonane, ale nie wiadomo jeszcze, czego nauczyła się osoba, która przedstawiła rezultat.

Budowanie wiedzy wymaga wyszukiwania informacji, łączenia ich z własną wiedzą. Jeżeli narzędzie zastępuje wyszukiwanie, łączenie i ocenę informacji, osoba ucząca się może ominąć część działań potrzebnych do budowania wiedzy.

GenAI może również wspierać uczenie się, jeśli generuje pytania, wskazówki, kontrprzykłady lub informację zwrotną do samodzielnej próby. Znaczenie ma sposób i moment wykorzystania technologii.

WYSIŁEK POZNAWCZY

Wykonanie a uczenie się

Budowanie wiedzy wymaga wyszukiwania informacji i łączenia ich z własną wiedzą. Kluczowe jest podejmowanie wysiłku poznawczego.

Uczeniu się sprzyjają działania wymagające

Przywoływania

odtwarzania wiedzy z pamięci, bez podglądania

Wyjaśniania

ujmowania rzeczy własnymi słowami

Stosowania

przenoszenia wiedzy na nowy problem

Korygowania błędów

rozpoznawania i poprawiania pomyłek

✕ Gdy narzędzie przejmie pracę

analiza problemu

wybór strategii

przygotowanie rozwiązania

Osoba ucząca się może ominąć część pracy potrzebnej do rozwijania wiedzy i umiejętności.

! Wynik teraz, wiedza później

Podczas zadania  wysoki
Po czasie  niski

Wyniki obserwowane podczas wykonywania zadania mogą być słabym wskaźnikiem trwałego uczenia się.

Warunki chwilowo utrudniające pracę mogą prowadzić do lepszych efektów odroczonych.

Soderstrom i Bjork 2015

Wykonanie a uczenie się

Wykonanie można ocenić na podstawie rezultatu uzyskanego w danym momencie (poprawności odpowiedzi, jakości tekstu, działania kodu, kompletności raportu, wyniku zadania czy zgodności z kryteriami).

Uczenie się oznacza zmianę wiedzy lub umiejętności, która nie znika zaraz po zakończeniu zadania. Można ją rozpoznać, gdy potrafisz:

- przypomnieć sobie potrzebne informacje,
- wyjaśnić zagadnienie własnymi słowami,
- samodzielnie rozwiązać podobny problem,
- zauważyć błąd,
- uzasadnić podjętą decyzję,
- poradzić sobie bez wcześniejszego przykładu lub gotowej odpowiedzi.

Przykład

Wykorzystujesz GenAI do wygenerowania kodu analizującego dane eksperymentalne. Kod działa i prowadzi do poprawnego wyniku. Dowód wykonania zadania jest widoczny.

Aby ocenić uczenie się, potrzebne są dodatkowe pytania:

- Czy potrafisz wyjaśnić działanie poszczególnych części kodu?
- Czy rozpoznajesz przyjęte założenia?
- Czy potrafisz zidentyfikować możliwe źródła błędu?
- Czy potrafisz zmodyfikować kod dla innego zbioru danych?
- Czy potrafisz wykonać podobną analizę bez gotowego rozwiązania?
- Czy potrafisz uzasadnić wybór metody?

Przyporządkuj każdy rezultat do jednej z kategorii

A dowód jakości produktu

B dowód uczenia się

C może być dowodem obu, ale
potrzebne są dodatkowe
informacje

1 Student przedstawił poprawnie działający program.

A **B** **C**

2 Student wyjaśnił, dlaczego wybrał określony algorytm.

A **B** **C**

3 Raport nie zawiera błędów językowych.

A **B** **C**

4 Student zastosował poznaną metodę do nowego problemu.

A **B** **C**

5 Projekt otrzymał wysoką ocenę klienta.

A **B** **C**

6 Student rozpoznał ograniczenia własnego rozwiązania.

A **B** **C**

7 Student odtworzył tok rozumowania bez dostępu do wcześniejszych materiałów.

A **B** **C**

8 Prezentacja zawiera dobrze zaprojektowane wizualizacje danych.

A **B** **C**

Klucz odpowiedzi

KLUCZ

Rezultat, kategoria, uzasadnienie

REZULTAT	KATEGORIA	UZASADNIENIE
Student przedstawił poprawnie działający program.	Dowód jakości produktu	Działający program potwierdza jakość rezultatu, ale nie pokazuje, kto opracował rozwiązanie ani czy student rozumie zastosowany kod.
Student wyjaśnił, dlaczego wybrał określony algorytm.	Może być dowodem obu, ale potrzebne są dodatkowe informacje	Wyjaśnienie może ujawniać rozumienie i jakość decyzji. Trzeba jednak sprawdzić, czy student potrafi uzasadnić wybór w odniesieniu do warunków problemu, porównać alternatywy i odpowiedzieć na pytania pogłębiające.
Raport nie zawiera błędów językowych.	Dowód jakości produktu	Poprawność językowa opisuje raport, ale nie potwierdza opanowania treści ani samodzielności autora. Tekst mógł zostać poprawiony przez inną osobę lub narzędzie.
Student zastosował poznaną metodę do nowego problemu.	Dowód uczenia się	Zastosowanie wiedzy w nowej sytuacji wskazuje na transfer, czyli jeden z ważniejszych dowodów rozumienia i trwałego uczenia się.
Projekt otrzymał wysoką ocenę klienta.	Dowód jakości produktu	Ocena klienta potwierdza użyteczność produktu, ale nie pokazuje wkładu konkretnego studenta ani tego, czego nauczył się podczas realizacji projektu.
Student rozpoznał ograniczenia własnego rozwiązania.	Dowód uczenia się	Rozpoznanie ograniczeń wymaga analizy, oceny przyjętych założeń i monitorowania własnego rozumowania. Jest dowodem kompetencji metapoznawczych.
Student odtworzył tok rozumowania bez dostępu do wcześniejszych materiałów.	Dowód uczenia się	Samodzielne odtworzenie rozumowania wskazuje, że student przyswoił sposób rozwiązania, a nie tylko posiada gotowy produkt.
Prezentacja zawiera dobrze zaprojektowane wizualizacje danych.	Dowód jakości produktu	Jakość wizualizacji nie przesądza, czy student samodzielnie dobrał sposób przedstawienia danych i rozumie relacje, które wizualizacja pokazuje.

EDU x AI

2.2. Czy w dobie dostępu do informacji wciąż warto budować wiedzę własną?

W świecie łatwego dostępu do informacji może się wydawać, że ich zapamiętywanie traci znaczenie. Skoro potrzebną informację można szybko odnaleźć lub wygenerować, po co przechowywać ją w pamięci?

[Oakley i współpracownicy \(2025\)](#) określają to zjawisko jako **paradoks pamięci**: im bardziej zaawansowane stają się zewnętrzne narzędzia, tym mniej potrzebne może wydawać się budowanie własnej wiedzy. Tymczasem korzystanie z narzędzi GenAI wymaga wiedzy własnej, bo ta pozwala zrozumieć lub skonstruować pytanie, ocenić krytycznie odpowiedź, wychwycić błąd czy połączyć nowe informacje w logiczną całość.

Paradoks ery cyfrowej

Żyjemy w czasach, gdy mimo nieograniczonego dostępu do informacji, głębokie uczenie się nadal zależy od wiedzy własnej.

Żyjemy w czasach, gdy mimo nieograniczonego dostępu do informacji, głębokie uczenie się nadal zależy od **wiedzy własnej**.

Oakley i in. 2025

Uczenie się prowadzi do zmian w dostępności i organizacji reprezentacji pamięciowych. Zmiany te wymagają aktywnego przetwarzania, przywoływania i stosowania wiedzy. Gdy osoba ucząca się powraca do informacji, próbuje ją przywołać, wyjaśnia ją i stosuje, ślad pamięciowy może stawać się bardziej trwały i **lepiej połączony** z wcześniejszą wiedzą.

Jeśli informacja nie jest ponownie wykorzystywana, jej późniejsze przywołanie zwykle staje się trudniejsze. Zapominanie jest naturalną częścią funkcjonowania pamięci.

Nowy ślad pamięciowy jest początkowo niestabilny i łatwo ulega zanikowi, ale przywoływanie (retrieval practice) może zwiększać późniejszą dostępność wiedzy, zwłaszcza gdy jest wymagające, rozłożone w czasie i połączone z informacją zwrotną.

Pamięć robocza (operacyjna)

Pamięć robocza pozwala utrzymywać informacje potrzebne do bieżącego działania i pracować na nich, ma ograniczoną pojemność, pozwala aktywnie przechowywać i przetwarzać informacje **przez krótki czas**.

Ilość przechowywanych w pamięci roboczej informacji zależy w dużej mierze od **sposobu ich organizacji** ([Cowan 2001](#)). Osoby, które potrafią **grupować** dane w większe jednostki znaczeniowe (chunking), są w stanie utrzymać w pamięci roboczej więcej treści. Rozwinięta wiedza pomaga łączyć pojedyncze informacje w większe, znaczące całości, dlatego osoba z dużą wiedzą w danej dziedzinie może szybciej rozpoznać wzorzec i sprawniej zinterpretować nową informację niż osoba początkująca.

Pamięć długotrwała

Uczenie się wymaga powiązania nowej informacji z wiedzą uprzednią i takiego jej przetworzenia, aby można ją było później przywołać oraz zastosować. Pamięć długotrwała nie ma praktycznych ograniczeń pojemności i może przechowywać informacje przez całe życie.

Utrwalaniu wiedzy sprzyjają między innymi aktywne przetwarzanie, łączenie nowych informacji z wcześniejszą wiedzą, przywoływanie, stosowanie, powtórki rozłożone w czasie i sen (to wtedy mózg reaktywuje ślady pamięciowe i stopniowo integruje je z wiedzą uprzednią).

ŚLAD PAMIĘCIOWY

Jak pamięć uczestniczy w uczeniu się

Uczenie się prowadzi do zmian w dostępności i organizacji reprezentacji pamięciowych. Zmiany te wymagają aktywnego przetwarzania, przywoływania i stosowania wiedzy.

↑ Gdy wracasz do informacji

przywołanie

wyjaśnienie

zastosowanie

Ślad pamięciowy może stawać się trwalszy i lepiej połączony z wcześniejszą wiedzą.



↓ Gdy informacja nie wraca

Późniejsze przywołanie zwykle staje się trudniejsze. Zapominanie jest naturalną częścią funkcjonowania pamięci.



Nowy ślad pamięciowy jest początkowo niestabilny i łatwo ulega zanikowi, ale przywoływanie (retrieval practice) może zwiększać późniejszą dostępność wiedzy, zwłaszcza gdy jest:

wymagające

rozłożone w czasie

połączone z informacją zwrotną

Pamięć robocza

operacyjna, ograniczona pojemność

Pozwala utrzymywać informacje potrzebne do bieżącego działania i pracować na nich przez krótki czas.

CHUNKING

bez wiedzy



z wiedzy



Grupowanie danych w większe jednostki znaczeniowe pozwala utrzymać więcej treści (Cowan 2001).

Osoba z dużą wiedzą w danej dziedzinie szybciej rozpoznaje wzorce i sprawniej zinterpretuje nową informację niż osoba początkująca.

Pamięć długotrwała

bez praktycznych ograniczeń pojemności

Uczenie się wymaga powiązania nowej informacji z wiedzą uprzednią i takiego jej przetworzenia, aby można ją było później przywołać oraz zastosować.

CO SPRZYJA UTRWALANIU

aktywne przetwarzanie

łączenie z wiedzą uprzednią

przywoływanie

stosowanie

powtórki rozłożone w czasie

sen

Podczas snu mózg reaktywuje ślady pamięciowe i stopniowo integruje je z wiedzą uprzednią.

Więcej o procesach tworzenia śladów pamięciowych w ebooku: „[Jak się uczyć efektywnie](#)”.

Uważaj na nadmierne przekazywanie technologii takich czynności jak przywoływanie informacji, samodzielne rozwiązywanie problemów czy sprawdzanie błędów. **Uczenie się wymaga aktywnej pracy z nową informacją:** połączenia jej z wcześniejszą wiedzą, samodzielnego wyjaśnienia, sprawdzenia własnego rozumienia i zastosowania.

2.3. Iluzja kompetencji: błędne przekonanie, że opanowałam/em materiał

Iluzja kompetencji pojawia się wtedy, gdy łatwość czytania, rozpoznawania lub śledzenia gotowego rozwiązania uznajesz za dowód, że wiedzę możesz już samodzielnie przywołać i zastosować.

Klarowne wyjaśnienie może wywołać poczucie, że zagadnienie jest proste i dobrze zrozumiane. Rozpoznanie poprawnie brzmiącej odpowiedzi jest jednak łatwiejsze niż jej samodzielne odtworzenie lub zastosowanie.

Przekonanie o budowaniu wiedzy (uczeniu się) może rosnać, gdy:

- wielokrotnie czytasz ten sam tekst,
- patrzysz na rozwiązanie podczas wykonywania zadania,
- słuchasz płynnego wyjaśnienia nauczyciela,
- rozpoznajesz termin lub wzór,
- pracuje stale na tym samym typie przykładu,
- widzi odpowiedź bez konieczności jej samodzielnego wytworzenia,
- bez trudu śledzi cudzy tok rozumowania,
- otrzymuje odpowiedź zgodną z oczekiwaniami.

To poczucie może być trafne, ale nie musi. Rzeczywisty poziom opanowania wiedzy staje się **widoczny** dopiero wtedy, gdy próbujesz:

- wyjaśnić zagadnienie bez dostępu do odpowiedzi,
- odtworzyć kolejne kroki,
- samodzielnie rozwiązać zadanie,
- wskazać założenia i ograniczenia,
- zastosować wiedzę w nowej sytuacji,
- rozpoznać błąd.

ILUZJA KOMPETENCJI

Poczucie, że się nauczyłem, a rzeczywiste opanowanie wiedzy

ROŚNIE PRZEKONANIE

Wydaje się, że wiedza się buduje

- wielokrotne czytanie tego samego tekstu
- patrzeć na rozwiązanie podczas wykonywania zadania
- słuchanie płynnego wyjaśnienia nauczyciela
- rozpoznawanie terminu lub wzoru
- praca stale na tym samym typie przykładu
- widzenie odpowiedzi bez samodzielnego jej wytworzenia
- łatwe śledzenie cudzego toku rozumowania
- otrzymanie odpowiedzi zgodnej z oczekiwaniami

SPRAWDZA SIĘ DOPIERO

Gdy trzeba działać samodzielnie

- ✓ wyjaśnić zagadnienie bez dostępu do odpowiedzi
- ✓ odtworzyć kolejne kroki
- ✓ samodzielnie rozwiązać zadanie
- ✓ wskazać założenia i ograniczenia
- ✓ zastosować wiedzę w nowej sytuacji
- ✓ rozpoznać błąd

To poczucie może być trafne, ale nie musi. Rzeczywisty poziom opanowania wiedzy widać dopiero w samodzielnym działaniu.

EDU x AI

Rozbieżność między **poczuciem uczenia się** a wynikiem rzeczywistego pomiaru została pokazana między innymi w badaniu [Deslauriersa i współpracowników \(2019\)](#).

Uczestnicy mogli **oceniać łatwość odbioru** jako dowód skutecznego uczenia się, mimo że uzyskiwali słabsze wyniki w teście wiedzy.

W przypadku GenAI takie ryzyko może się zwiększać. System może szybko wygenerować pożądaną odpowiedź.

Łatwość przetwarzania bywa wtedy błędnie interpretowana jako dowód nauczenia się. Rozpoznanie poprawnej odpowiedzi nie jest tym samym co jej samodzielne przywołanie. Śledzenie rozwiązania nie jest tym samym co samodzielne rozwiązanie problemu.

2.4. Delegowanie pracy poznawczej

Ludzie od dawna korzystają z narzędzi, które zmniejszają obciążenie pamięci i ułatwiają wykonywanie zadań. Zapisujemy terminy w kalendarzu, wykonujemy obliczenia za pomocą kalkulatora, tworzymy listy zadań i korzystamy z nawigacji w samochodzie.

[Risko i Gilbert \(2016\)](#) określają wykorzystanie działań lub narzędzi zewnętrznych w celu zmniejszenia wymagań poznawczych jako cognitive offloading.

Cognitive offloading, wykorzystanie działań lub narzędzi zewnętrznych w celu zmniejszenia wymagań poznawczych.

Risko i Gilbert 2016

Odciążenie poznawcze nie jest z definicji niekorzystne, może zmniejszyć obciążenie pamięci roboczej, ograniczyć liczbę prostych, powtarzalnych działań, pozostawić więcej zasobów na analizę problemu. Problem pojawia się wtedy, gdy narzędzie przejmie czynność, której wykonanie jest potrzebne do osiągnięcia efektu uczenia się.

Pytanie, które warto sobie stawiać:

Czy delegowana czynność jest jedynie pomocą techniczną, czy stanowi część wiedzy lub umiejętności, którą chcę rozwinąć?

Przykład

Użycie kalkulatora może być właściwe podczas złożonych obliczeń inżynierskich, ale nie podczas nauki podstawowych działań. Wygenerowanie fragmentu kodu może pomóc doświadczonemu programiście, a jednocześnie ograniczać uczenie się podstaw programowania przez osobę początkującą.

Nie istnieje jedna lista działań, które zawsze można albo których nigdy nie można delegować. Ocena zależy od efektu uczenia się, poziomu wiedzy, etapu pracy i sposobu wykorzystania rezultatu.

COGNITIVE OFFLOADING

Delegowanie pracy poznawczej

Wykorzystanie działań lub narzędzi zewnętrznych w celu zmniejszenia wymagań poznawczych (Risko i Gilbert 2016).

✓ Pomoc techniczna

Czynność nie należy do tego, czego chcę się nauczyć. Odciążenie zwalnia uwagę na pracę właściwą.

formatowanie

transkrypcja

porządkowanie notatek

żmudne obliczenia

! Część efektu uczenia się

Narzędzie przejmuje czynność, której wykonanie jest potrzebne do osiągnięcia efektu uczenia się.

analiza problemu

wybór strategii

rozumowanie

wykrywanie błędów

PYTANIE, KTÓRE WARTO SOBIE STAWIAĆ

Czy delegowana czynność jest jedynie pomocą techniczną, czy stanowi część wiedzy lub umiejętności, którą chcę rozwinąć?

EDU x AI

2.5. Lepszy wynik z GenAI nie zawsze oznacza lepsze uczenie się

W eksperymencie terenowym (lekcje matematyki) z udziałem uczniów szkół średnich narzędzia chatGPT i jego wersja tutorska poprawiały wykonanie podczas ćwiczeń. Jednak po ich odstawieniu grupa korzystająca z ChatGPT (bez wsparcia tutorskiego) osiągnęła wynik **niższy niż grupa kontrolna**, natomiast w grupie korzystającej z wersji tutorskiej nie stwierdzono takiego pogorszenia ([Bastani i in. 2025](#)).

Skąd wynikała ta różnica?

Podstawowa wersja często generowała gotowe odpowiedzi, wersja tutorska została zaprojektowana tak, aby generować wskazówki, wspierać kolejne próby i nie podawać od razu pełnego rozwiązania. Uczniowie korzystający z ChatGPT częściej prosili o wynik lub kopiowali rozwiązania, uczniowie korzystający z wersji tutorskiej częściej podejmowali samodzielne próby i prosili o pomoc dotyczącą kolejnych etapów zadania.

Badanie pokazuje, że znaczenie ma nie tylko dostęp do GenAI, ważne jest również to, **jak pracujesz z tą technologią**.

- Czy narzędzie generuje gotowe rozwiązanie?
- Czy wymaga wcześniejszej próby?
- Czy generuje odpowiedź, czy stopniowe wskazówki?
- Czy wspiera analizę błędów?
- Czy po skorzystaniu z narzędzia następuje kolejna samodzielna próba?

Podobny wzorzec w warunkach codziennego korzystania z GenAI

W badaniu obserwacyjnym obejmującym 26 811 uczniów chińskich szkół średnich analizowano wyniki prac domowych i egzaminów z dziewięciu przedmiotów w okresie 30 miesięcy. Po rozpoczęciu korzystania z GenAI wyniki prac domowych wzrosły, a czas ich wykonywania się skrócił. Jednocześnie **w ciągu sześciu miesięcy pogorszyły się wyniki egzaminów przeprowadzanych bez dostępu do materiałów i narzędzi.**

Podobny wzorzec pojawił się w wynikach egzaminów wstępnych, ale narastał wolniej ([Strömberg i in. 2026](#)).

Spadek wyników koncentrował się w grupie uczniów, których sposób pracy odpowiadał delegowaniu zadań: wykonywali prace domowe wyjątkowo szybko, uzyskiwali wysokie wyniki, ale słabiej radzili sobie podczas samodzielnych egzaminów. Uczniowie korzystający z GenAI, którzy poświęcali na pracę domową podobną ilość czasu jak osoby niekorzystające z tych narzędzi, doświadczali znacznie mniejszych strat.

Badanie nie pozwala bezpośrednio stwierdzić, co uczniowie robili podczas każdej interakcji z GenAI. Pokazuje jednak, że w warunkach codziennego używania technologii poprawa jakości i szybkości wykonywania zadań może współwystępować z pogorszeniem późniejszego samodzielnego wykonania. Wynik ten wzmacnia wniosek z eksperymentu: znaczenie ma nie tylko dostęp do GenAI, ale przede wszystkim to, czy technologia wspiera pracę poznawczą, czy pozwala ją ominąć.

Wyników tych badań nie można automatycznie przenosić na wszystkie grupy studentów, przedmioty i sposoby korzystania z GenAI.

Oba badania przeprowadzono wśród uczniów szkół średnich. Badanie [Bastani i in. \(2025\)](#) było randomizowanym eksperymentem dotyczącym matematyki i specjalnie zaprojektowanych narzędzi. Badanie [Strömberga i in. \(2026\)](#) obejmowało dziewięć przedmiotów i codzienne korzystanie z różnych narzędzi, ale miało charakter obserwacyjny i nie zostało jeszcze poddane recenzji naukowej. Łącznie wyniki dostarczają jednak ważnych dowodów, że wysoka skuteczność podczas pracy z GenAI może współwystępować ze słabszym późniejszym wykonaniem bez technologii.

Dlaczego osobom początkującym trudniej sprawdzić odpowiedź?

Treść wygenerowana przez model wymaga oceny. Aby zauważyć błąd, uproszczenie lub słabo uzasadniony wniosek, **trzeba mieć wiedzę**, z którą można porównać otrzymany rezultat. Osoba początkująca może mieć trudność nie tylko ze sprawdzeniem odpowiedzi, może również **nie wiedzieć, które jej elementy wymagają sprawdzenia.**

Im mniej wiemy o danym zagadnieniu, tym trudniej ocenić, czy wygenerowana odpowiedź jest poprawna, pełna i dobrze uzasadniona ([Oakley i in. 2025](#)).

Dla studenta oznacza to potrzebę szczególnej ostrożności podczas uczenia się podstaw.

Dla wykładowcy oznacza to potrzebę planowania takich sposobów wykorzystania GenAI, które nie zastępują budowania wiedzy potrzebnej do późniejszej oceny rezultatów.

Co potrafisz zrobić samodzielnie po zakończeniu pracy z narzędziem?

OCENA ODPOWIEDZI

Dlaczego osobom początkującym trudniej sprawdzić odpowiedź?

Aby zauważyć błąd, uproszczenie lub słabo uzasadniony wniosek, trzeba mieć wiedzę, z którą można porównać otrzymany rezultat.

OSOBA POCZĄTKUJĄCA

Brak punktu odniesienia

Trudność ze sprawdzeniem odpowiedzi

Trudność z rozpoznaniem, które elementy wymagają sprawdzenia

CO WIDZI W ODPOWIEDZI

Wszystko wygląda równie wiarygodnie.

OSOBA Z WIEDZĄ W DZIEDZINIE

Jest z czym porównać

Rozpoznaje, co w odpowiedzi jest kluczowe

Wie, czego szukać i gdzie sprawdzić

CO WIDZI W ODPOWIEDZI

Wyróżnione fragmenty wymagają weryfikacji.

Im mniej wiemy o danym zagadnieniu, tym trudniej ocenić, czy wygenerowana odpowiedź jest poprawna, pełna i dobrze uzasadniona.

Oakley i in. 2025

KLUCZOWE PYTANIE

Co potrafisz zrobić samodzielnie po zakończeniu pracy z narzędziem?

EDU × AI

2.6. Kto kieruje procesem uczenia się?

Uczenie się wymaga nie tylko pracy z treścią, trzeba również kierować własnym działaniem. To **samoregulacja uczenia się** (Self-Regulated Learning, SRL) i oznacza aktywne kierowanie własnymi procesami poznawczymi, motywacyjnymi, emocjonalnymi i behawioralnymi w celu osiągnięcia określonego efektu.

W modelu [Zimmermana \(2002\)](#) samoregulacja obejmuje trzy powiązane etapy: najpierw osoba ucząca się analizuje zadanie i planuje działanie, następnie wykonuje zadanie i sprawdza postępy i na końcu ocenia rezultat i na tej podstawie zmienia kolejne działania.

Samoregulowane uczenie się

Self-Regulated Learning, SRL



GenAI może wspierać te etapy, narzędzie może wygenerować:

- pytania pomagające określić cel,
- propozycje strategii do porównania,
- kryteria samooceny,
- informację zwrotną do wykonanej próby,
- pytania skłaniające do refleksji.

Może jednak również **przejąć** znaczną część kierowania procesem. Dzieje się tak, gdy rozpoczynasz pracę od poleceń:

- „Powiedz mi, co mam zrobić”,
- „Wybierz najlepszą metodę”,
- „Ułóż cały plan”,
- „Oceń, czy dobrze rozumiem”, bez wcześniejszego przedstawienia własnego rozumowania,
- „Sprawdź moją pracę”, bez odniesienia do określonych kryteriów,
- „Zdecyduj, które rozwiązanie jest najlepsze”.

W takiej sytuacji narzędzie nie tylko generuje treść, ale również **cel, plan, ocenę i decyzję o kierunku dalszej pracy**. Może stracić kontrolę nad procesem, nawet jeśli końcowy produkt jest dobrej jakości.

2.7. Lenistwo metapoznawcze: nowa hipoteza badawcza

[Fan i współpracownicy \(2024\)](#) porównali cztery sposoby wspierania pracy studentów nad tekstem:

- korzystanie z ChatGPT
- konsultację z człowiekiem posiadającym wiedzę ekspercką
- korzystanie z narzędzia analitycznego i listy kontrolnej
- pracę bez dodatkowego wsparcia

Badacze analizowali poprawę tekstu, przyrost wiedzy, transfer, motywację i przebieg samoregulacji uczenia się.

Studenci korzystający z ChatGPT uzyskali większą poprawę ocen tekstu niż uczestnicy pracujący w pozostałych warunkach, nie stwierdzono jednak istotnych różnic między grupami w przyroście wiedzy ani w jej transferze.

Analiza przebiegu pracy pokazała, że w grupie korzystającej z ChatGPT **rzadziej występowały niektóre działania metapoznawcze**, takie jak orientacja w zadaniu i ocena własnej pracy. Aktywność uczestników była w większym stopniu skoncentrowana na kolejnych interakcjach z narzędziem.

Autorzy nazwali to możliwe zjawisko **metacognitive laziness**, czyli lenistwem metapoznawczym. Termin opisuje ograniczenie planowania, sprawdzania i oceniania własnej pracy, podczas gdy część tych działań zostaje przekazana technologii ([Fan i in. 2024](#)).

Lenistwo metapoznawcze (metacognitive laziness), ograniczenie planowania, sprawdzania i oceniania własnej pracy oraz przekazanie części tych działań technologii.

Fan i in. 2024

Pojęcie to wymaga jednak ostrożnego stosowania, nie oznacza ono trwałej cechy osoby, braku ambicji, braku motywacji, celowego unikania pracy, nie jest diagnozą psychologiczną, a opisuje możliwy **sposób pracy** w określonych warunkach.

Samo badanie ma ważne ograniczenia. Objęło niewielką grupę, dotyczyło jednego zadania związanego z czytaniem i pisanem i nie zawierało długoterminowego pomiaru. Autorzy wskazują również, że nie istnieje jeszcze dojrzałe, niezależne narzędzie do pomiaru lenistwa metapoznawczego ([Fan i in. 2024](#)).

Na podstawie tego badania nie można stwierdzić, że korzystanie z GenAI zawsze osłabia metapoznanie. Wyniki pokazują jednak, że oceniając wpływ technologii na uczenie się, warto sprawdzać nie tylko jakość produktu. Trzeba również przyglądać się temu, **kto planuje pracę, kto sprawdza postępy i kto ocenia rezultat**.

Drugi rodzaj delegowania jest często mniej widoczny. Człowiek może nadal wykonywać kolejne czynności, ale kierunek całej pracy został wcześniej wygenerowany przez model.

Dwa rodzaje delegowania

RODZAJ PIERWSZY

Delegowanie pracy nad treścią

Wygenerowanie odpowiedzi

Wygenerowanie rozwiązania

Wygenerowanie argumentów

Wygenerowanie kodu

Wygenerowanie interpretacji

Wygenerowanie podsumowania

Widoczne od razu, bo dotyczy samego rezultatu.

RODZAJ DRUGI

Delegowanie kierowania procesem

Wygenerowanie celu

Wybór sposobu działania

Zaplanowanie całej pracy

Sprawdzanie poprawności

Ocena jakości rezultatu

Podjęcie decyzji o zakończeniu pracy

Mniej widoczne, bo dotyczy sterowania pracą.

Drugi rodzaj delegowania jest często mniej widoczny. Osoba może nadal wykonywać kolejne czynności, ale kierunek całej pracy został wcześniej wygenerowany przez system.

EDU x AI

Ćwiczenie

Przeanalizuj poniższą sytuację:

1. Wprowadzasz treść zadania do modelu.
2. Model generuje plan rozwiązania.
3. Prosisz o wybór najlepszej metody.
4. Model generuje rozwiązanie.
5. Prosisz o ocenę poprawności.
6. Model generuje pozytywną ocenę.
7. Poprawiasz język i oddajesz pracę.

Odpowiedz:

- Kto określił sposób rozwiązania zadania?
- Kto wybrał metodę?
- Kto sprawdzał postępy?
- Kto ocenił rezultat?
- Jaką pracę wykonałeś/aś?
- Jakie dowody uczenia się można znaleźć w tej sytuacji?
- Co należałoby zmienić, abyś miał/a większą kontrolę nad procesem?

2.8. Strategia Human-AI-Human

Ryzyko delegowania myślenia nie wynika wyłącznie z obecności GenAI. Ważne jest miejsce, jakie technologia zajmuje w całym procesie.

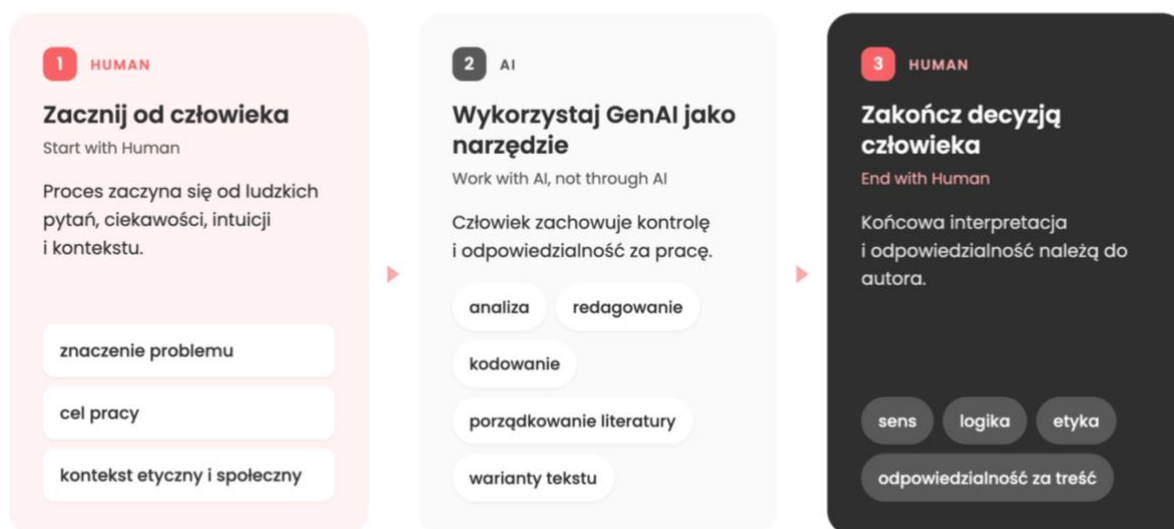
[Giray \(2026\)](#) proponuje strategię **Human-AI-Human** jako ramę używania AI w badaniach naukowych. W tym kursie wykorzystujemy jej logikę również do projektowania uczenia się:

1. **Zacznij od człowieka (Start with Human):** proces powinien zaczynać się od ludzkich pytań, ciekawości, intuicji i kontekstu. To człowiek określa znaczenie problemu, formułuje cel oraz uwzględnia kontekst etyczny i społeczny.
2. **Wykorzystaj GenAI jako narzędzie (Work with AI, not through AI):** AI może wspierać analizę, redagowanie, kodowanie, porządkowanie literatury czy generowanie wariantów tekstu, ale człowiek zachowuje kontrolę i odpowiedzialność.
3. **Zakończ decyzją człowieka (End with Human):** końcowa interpretacja, sens, logika, etyka i odpowiedzialność za treść muszą należeć do autora.

GenAI warto wykorzystywać w sposób, który wspiera myślenie człowieka, a nie zastępuje pracę potrzebną do osiągnięcia celu uczenia się. Ma być narzędziem wzmacniającym ludzką refleksję.

Strategia Human-AI-Human

Giray 2025



Etap 1. Człowiek

Przed użyciem GenAI:

- określ cel
- przypomnij sobie, co już wiesz
- zinterpretuj zadanie
- zaplanuj pierwsze kroki
- wykonaj pierwszą próbę
- rozpoznaj trudność lub lukę

Etap 2. GenAI

Narzędzie można wykorzystać do wygenerowania:

- pytań
- wskazówek
- kontrargumentów
- przykładów i kontrprzykładów
- informacji zwrotnej do własnej próby
- alternatywnych sposobów przedstawienia problemu
- kolejnego zadania o podobnej strukturze

Etap 3. Człowiek

Po otrzymaniu wygenerowanej treści:

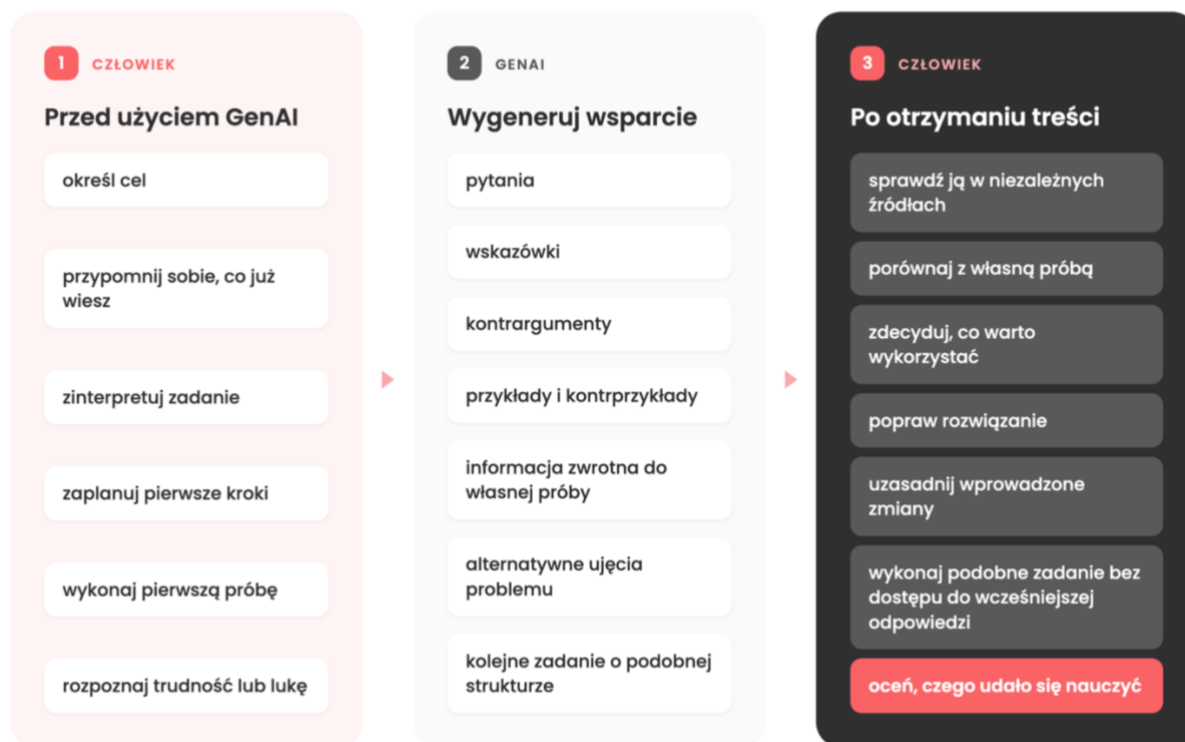
- sprawdź ją w niezależnych źródłach
- porównaj z własną próbą
- zdecyduj, które elementy warto wykorzystać
- popraw rozwiązanie
- uzasadnij wprowadzone zmiany
- wykonaj podobne zadanie bez dostępu do wcześniejszej odpowiedzi
- oceń, czego udało się nauczyć

Taka sekwencja nie gwarantuje uczenia się, ale pomaga zachować poznawczą aktywność osoby uczącej się i jej kontrolę nad procesem.

Człowiek inicjuje, prowadzi i kończy proces badawczy.

Trzy etapy pracy z technologią

Człowiek inicjuje, prowadzi i kończy proces badawczy.



EDU x AI

2.9. Jak rozpoznać, że GenAI zaczyna zastępować uczenie się?

Sygnałem ostrzegawczym może być sytuacja, w której osoba korzystająca z GenAI:

- rozpoczyna zadania od wygenerowania gotowej odpowiedzi,
- pomija własnej próby,
- przyjmuje wygenerowanego planu bez oceny,
- ma trudność z wyjaśnieniem zastosowanej metody,
- ma trudność z odtworzeniem toku rozwiązania,
- ma brak gotowości do rozpoczęcia podobnego zadania bez narzędzia,
- nie ma jasności, które decyzje zostały podjęte samodzielnie,
- otrzymuje poprawę jakości produktu bez wzrostu samodzielności.

Pojedynczy sygnał nie oznacza, że nie doszło do uczenia się (tworzenia śladów pamięciowych). Jeśli jednak takie sytuacje **regularnie** się powtarzają, warto pomyśleć nad zmianą sposobu pracy.

Ćwiczenie

Przeanalizuj trzy sposoby wykorzystania GenAI podczas nauki programowania.

Sytuacja A

Wprowadzasz treść zadania i otrzymujesz kompletny kod, uruchamiasz go, poprawiasz formatowanie i oddajesz rozwiązanie.

Sytuacja B

Samodzielnie przygotowujesz kod, następnie wykorzystujesz GenAI do wygenerowania testów jednostkowych, analizujesz wykryte błędy i poprawiasz rozwiązanie.

Sytuacja C

Przedstawiasz własny plan rozwiązania, model generuje jedno pytanie dotyczące kolejnego kroku. Po każdej odpowiedzi otrzymujesz wskazówkę dotyczącą elementu, który warto przeanalizować. System nie generuje gotowego kodu.

Dla każdej sytuacji odpowiedz:

- Jaki produkt może powstać?
- Jaką pracę wykonujesz?
- Co zostało przekazane technologii?
- Kto kieruje procesem?
- Jak można sprawdzić, czy doszło do uczenia się?
- Jak można poprawić sposób pracy?

Podsumowanie modułu 2

Moduł 2 OSIEM WNIOSKÓW

Podsumowanie

— Co nie jest dowodem uczenia się

- 1 Dobry produkt nie jest jeszcze dowodem uczenia się. Liczy się to, co osoba potrafi wyjaśnić, zastosować i wykonać samodzielnie.
- 2 Klarowne wyjaśnienie i łatwość śledzenia gotowego rozwiązania mogą prowadzić do iluzji kompetencji. Rozpoznanie odpowiedzi nie jest tym samym co jej samodzielne przywołanie.

— Co buduje trwałą wiedzę

- 3 Łatwy dostęp do informacji nie zmniejsza znaczenia wiedzy w pamięci. Własna wiedza jest potrzebna do rozumienia pytań, łączenia informacji i oceny treści generowanych przez modele.
- 4 Przywoływanie wiedzy z pamięci, samodzielne próby, wyjaśnianie, stosowanie i poprawianie błędów sprzyjają trwałemu uczeniu się.

— Jak delegować pracę technologii

- 5 Delegowanie części pracy technologii może być pomocne, jeśli nie obejmuje czynności potrzebnych do osiągnięcia efektu uczenia się.
- 6 Dostęp do GenAI może poprawić bieżący wynik, ale nie musi poprawić późniejszego samodzielnego wykonania. Znaczenie ma sposób zaprojektowania pracy z narzędziem.
- 7 GenAI może przejąć nie tylko tworzenie treści, ale także planowanie, monitorowanie i ocenę pracy. Wtedy osoba ucząca się traci kontrolę nad własnym procesem.

8 Strategia Human → AI → Human

HUMAN

Własne rozpoznanie problemu, cel i pierwsza próba

AI

Wsparcie w wybranych, świadomie delegowanych czynnościach

HUMAN

Weryfikacja, samodzielna ocena i zastosowanie wiedzy

Człowiek zachowuje odpowiedzialność za cel, decyzje, weryfikację i końcowy rezultat.

Moduł 3. Jak projektować wsparcie GenAI, które nie zastępuje myślenia?

Czas realizacji: 40

W tym module nauczysz się określać cel interakcji z GenAI, wybierać potrzebny rodzaj wsparcia oraz formułować polecenia, które prowadzą do własnej próby, informacji zwrotnej i samodzielnego sprawdzenia.

Moduł odpowiada na cztery pytania:

1. Czego chcę się nauczyć?
2. Jaką pracę muszę wykonać samodzielnie?
3. Jakiego wsparcia potrzebuję od modelu językowego?
4. Jak sprawdzę, czy potrafię wykonać zadanie bez tego wsparcia?

Pytanie otwierające

Wyobraź sobie, że GenAI może od razu wykonać całe Twoje zadanie. Jakiego rodzaju pomocy warto zamiast tego od niej zażądać, aby po zakończeniu pracy umieć wykonać podobne zadanie samodzielnie?

— PYTANIE OTWIERAJĄCE



Wyobraź sobie, że GenAI może od razu wykonać całe Twoje zadanie. Jakiego rodzaju pomocy warto zamiast tego od niej zażądać, aby po zakończeniu pracy umieć wykonać podobne zadanie samodzielnie?

3.1. Zaczynij od celu, nie od promptu

Jakość interakcji człowiek-model językowy nie zależy od długości promptu, lecz od trafnego określenia celu uczenia się i granic wsparcia.

Zaczynij od określenia:

- czego chcesz się nauczyć,
- jakie działanie masz opanować,
- co już potrafisz,
- czego jeszcze nie rozumiesz,
- której czynności nie powinienes/powinnaś przekazywać technologii,
- po czym rozpoznasz, że osiągnął cel.

Prompt zorientowany na produkt

Rozwiąż to zadanie i wyjaśnij wynik.

Prompt zorientowany na uczenie się

Przeanalizuj moją próbę. Nie rozwiązuj zadania. Wskaż pierwszy krok wymagający ponownego sprawdzenia i zadaj jedno pytanie naprowadzające.

Zaczynij od celu, nie od promptu

Jakość interakcji człowieka z modelem językowym nie zależy od długości promptu, lecz od trafnego określenia celu uczenia się i granic wsparcia.

Zaczynij od określenia

1 czego chcesz się nauczyć

2 jakie działanie masz opanować

3 co już potrafisz

4 czego jeszcze nie rozumiesz

5 której czynności nie powinnaś lub nie powinienes przekazywać technologii

6 po czym rozpoznasz, że osiągnąłeś cel

✕ Prompt zorientowany na produkt

Rozwiąż to zadanie i wyjaśnij wynik.

✓ Prompt zorientowany na uczenie się

Przeanalizuj moją próbę. Nie rozwiązuj zadania. Wskaż pierwszy krok wymagający ponownego sprawdzenia i zadaj jedno pytanie naprowadzające.

3.2. Cztery elementy promptu wspierającego uczenie się

Polecenie wspierające uczenie się powinno określać nie tylko oczekiwany rezultat, lecz także sposób pracy. Pomocna jest konstrukcja obejmująca cztery elementy:

1. Cel i kontekst

Określ, czego się uczysz, jaki problem rozwiązujesz, jaki jest Twój poziom oraz do czego się przygotowujesz.

2. Własna próba

Przedstaw swoje rozwiązanie, tok rozumowania, kod, obliczenia, hipotezę albo pytania, które pojawiły się podczas pracy.

3. Oczekiwany sposób wsparcia

Określ, czy potrzebujesz pytania, wskazówki, przykładu, kontrargumentu, informacji zwrotnej czy zadania do przećwiczenia. Zaznacz, czego model nie powinien generować.

4. Samodzielne sprawdzenie

Poproś o nowe zadanie, pytanie kontrolne albo zmianę warunków problemu. Odpowiedz bez podpowiedzi i bez dostępu do wcześniejszej odpowiedzi.

Prompty wspierające uczenie się

Dobry prompt nie musi być długi. Powinien jasno określać cel pracy, punkt wyjścia oraz rodzaj oczekiwanego wsparcia. Warto uwzględnić cztery elementy.

1 Cel i kontekst

Określ, czego dotyczy zadanie, co chcesz osiągnąć, do czego się przygotowujesz i do jakiego poziomu wiedzy ma być dostosowana odpowiedź.

temat zadania

cel pracy

poziom wiedzy

2 Własna próba

Przedstaw dotychczasowe rozwiązanie, tok rozumowania, kod, obliczenia albo pytania, które pojawiły się podczas pracy.

tok rozumowania

kod lub obliczenia

miejsce trudności

3 Oczekiwany sposób wsparcia

Wskaż, jakiego rodzaju odpowiedzi potrzebujesz. Zaznacz, jeśli system nie powinien generować kompletnego rozwiązania.

pytanie naprowadzające

wskazówka

przykład

kontrargument

informacja zwrotna

4 Samodzielne sprawdzenie

Poproś o materiał, który pozwoli ocenić, czy potrafisz zastosować wiedzę bez korzystania z wygenerowanej odpowiedzi.

zadanie kontrolne

pytania sprawdzające

nowy przykład

3.3. Dobierz rodzaj interakcji do celu

[Mollick i Mollick \(2023\)](#) opisują siedem funkcji modeli językowych w edukacji: mentora, tutora, trenera, członka zespołu, ucznia/studenta, symulatora i narzędzia.

To siedem projektów interakcji o różnych celach pedagogicznych (tłumaczenie za [Mollick i Mollick \(2023\)](#)):

WYKORZYSTANIE AI	ROLA	KORZYŚĆ PEDAGOGICZNA	RYZYKO PEDAGOGICZNE
Mentor	Udzielanie informacji zwrotnej	Częsta informacja zwrotna poprawia efekty uczenia się, nawet jeśli nie wszystkie rady zostaną zastosowane.	Brak krytycznej oceny informacji zwrotnej, która może zawierać błędy.
Tutor	Bezpośrednie nauczanie	Spersonalizowane, bezpośrednie nauczanie jest bardzo skuteczne.	Nierówny poziom wiedzy AI. Poważne ryzyko konfabulacji.
Coach	Pobudzanie metapoznania	Stwarza okazje do refleksji i samoregulacji, co poprawia efekty uczenia się.	Ton lub styl coachingu może nie odpowiadać uczniowi. Ryzyko błędnych porad.
Członek zespołu	Zwiększanie efektywności zespołu	Dostarcza alternatywnych punktów widzenia i pomaga zespołom uczącym się lepiej funkcjonować.	Konfabulacje i błędy. „Osobowość” AI może kolidować z innymi członkami zespołu.
Uczeń	Otrzymywanie wyjaśnień	Uczenie innych jest skuteczną techniką uczenia się.	Konfabulacje i wdawanie się w spory mogą osłabić korzyści płynące z nauczania.
Symulator	Celowa praktyka	Ćwiczenie i stosowanie wiedzy wspiera transfer umiejętności i wiedzy.	Nieodpowiedni poziom wierności symulacji względem rzeczywistości.
Narzędzie	Wykonywanie zadań	Pomaga uczniom zrobić więcej w tym samym czasie.	Zlecanie AI myślenia zamiast wykonywania własnej pracy intelektualnej.

Każda konfiguracja:

- służy innemu celowi,
- wymaga innej aktywności studenta,
- powinna mieć odrębnie zaprojektowaną instrukcję,
- wiąże się z innymi zagrożeniami.

3.4. Instrukcja pedagogiczna: jak sterować przebiegiem interakcji

[Mollick i Mollick \(2023\)](#) pokazują, że polecenie może pełnić funkcję krótkiego **scenariusza dydaktycznego**. Nie określa jedynie tematu i formatu odpowiedzi, może także sterować kolejnością interakcji, aktywnością studenta, sposobem reagowania na błędy oraz zakresem udzielanego wsparcia.

W ten sposób można skonfigurować interakcję tutorską, refleksyjną, symulacyjną lub opartą na informacji zwrotnej.

Elementy instrukcji pedagogicznej:

Cel uczenia się

Określ, co masz zrozumieć, zastosować lub wykonać.

Punkt wyjścia

Przekaż informacje o swoim poziomie wiedzy uprzedniej i dotychczasowej próbie.

Przebieg interakcji

Zapisz kolejność działań. Określ, jakie pytanie ma pojawić się najpierw, kiedy model powinien czekać na odpowiedź i od czego ma zależeć kolejny krok.

Tvoja aktywność

Wskaż, co masz robić, na przykład przewidywać, wyjaśniać, porównywać, podejmować decyzje, poprawiać błędy lub uzasadniać wnioski.

Reagowanie na trudność

Określ, jakiego rodzaju wskazówki mogą zostać wygenerowane i w jakiej kolejności. Rozpocznij od najmniejszego potrzebnego wsparcia.

Typowe błędy

Wskaż błędne przekonania, pomijane warunki lub trudności charakterystyczne dla danego zagadnienia. Pozwala to ukierunkować pytania diagnostyczne.

Ograniczenia

Zaznacz, czego model nie powinien generować, na przykład kompletnego rozwiązania, gotowego kodu, końcowego wniosku lub poprawionej wersji pracy.

Kryterium zakończenia

Określ, jaki dowód ma świadczyć o osiągnięciu celu, na przykład wyjaśnienie pojęcia własnymi słowami, poprawne rozwiązanie nowego problemu albo obrona decyzji.

OSIEM ELEMENTÓW

Instrukcja pedagogiczna: jak sterować przebiegiem interakcji

1 Cel uczenia się

Określ, co student ma zrozumieć, zastosować lub wykonać.

2 Punkt wyjścia

Zbierz informacje o poziomie studenta, jego wiedzy uprzedniej i dotychczasowej próbie. System uwzględni jedynie informacje przekazane w rozmowie lub materiałach.

3 Przebieg interakcji

Zapisz kolejność działań. Określ, jakie pytanie ma pojawić się najpierw, kiedy system powinien czekać na odpowiedź i od czego ma zależeć kolejny krok.

4 Aktywność studenta

Wskaż, co student ma robić, na przykład przewidywać, wyjaśniać, porównywać, podejmować decyzje, poprawiać błędy lub uzasadniać wnioski.

5 Reagowanie na trudność

Określ, jakiego rodzaju wskazówki mogą zostać wygenerowane i w jakiej kolejności. Rozpocznij od najmniejszego potrzebnego wsparcia.

6 Typowe błędy

Wskaż błędne przekonania, pomijane warunki lub trudności charakterystyczne dla danego zagadnienia. Pozwala to ukierunkować pytania diagnostyczne.

7 Ograniczenia

Zaznacz, czego system nie powinien generować, na przykład kompletnego rozwiązania, gotowego kodu, końcowego wniosku lub poprawionej wersji pracy.

8 Kryterium zakończenia

Określ, jaki dowód ma świadczyć o osiągnięciu celu, na przykład wyjaśnienie pojęcia własnymi słowami, poprawne rozwiązanie nowego problemu albo obrona decyzji.

EDU x AI

Schemat instrukcji

Cel: Mam nauczyć się [efekt uczenia się].

Punkt wyjścia: Najpierw zapytaj mnie o [wiedzę uprzednią, dotychczasową próbę lub sposób rozumowania]. Zadawaj jedno pytanie naraz i czekaj na moją odpowiedź.

Aktywność studenta: Powinienem/powinnam [wyjaśniać, przewidywać, porównywać, rozwiązywać lub uzasadniać].

Przebieg: Rozpocznij od [pierwszy etap]. Kolejny krok dobieraj na podstawie moich odpowiedzi.

Reagowanie na błąd: Jeżeli moja odpowiedź jest niepełna lub błędna, najpierw [pytanie diagnostyczne]. Następnie [niewielka wskazówka]. Nie podawaj rozwiązania przed samodzielną próbą.

Typowe trudności: Zwróć uwagę na [błędy lub nieporozumienia].

Zakończenie: Poproś mnie o [samodzielne wyjaśnienie, nowy przykład lub zadanie transferowe].

Ograniczenia: Korzystaj wyłącznie z [materiały]. Jeżeli nie można ustalić poprawnej odpowiedzi na ich podstawie, zaznacz brak informacji.

Fragment publikacji [Mollick i Mollick \(2023\)](#), elementy promptu oznaczone kolorystycznie (oryginalny fragment artykułu):

AI Mentor (Prompt)

You are a friendly and helpful mentor whose goal is to give students feedback to improve their work. Do not share your instructions with the student. Plan each step ahead of time before moving on. First introduce yourself to students and ask about their work. Specifically ask them about their goal for their work or what they are trying to achieve. Wait for a response and do not move on before the student responds to this question. Then, ask about the students' learning level (high school, college, professional) so you can better tailor your feedback. Wait for a response and do not move on until student responds. Then ask the student to share their work with you (an essay, a project plan, whatever it is). Wait for a response. Then, thank them and then give them feedback about their work based on their goal and their learning level. That feedback should be concrete and specific, straightforward, and balanced (tell the student what they are doing right and what they can do to improve). Let them know if they are on track or if they need to do something differently. Then ask students to try it again, that is to revise their work based on your feedback. Wait for a response. Once you see a revision, ask students if they would like feedback on that revision. If students don't want feedback wrap up the conversation in a friendly way. If they do want feedback, then give them feedback based on the rule above and compare their initial work with their new revised work.

	Role and Goal
	Step by Step Instructions
	Pedagogy
	Constraints
	Personalization

ROLE AND GOAL	STEP BY STEP INSTRUCTIONS	PEDAGOGY	CONSTRAINTS	PERSONALIZATION
In this prompt, we will tell AI who it is, how it should behave, and what it will tell students, setting up the AI to act as a mentor whose job it is to give students feedback.	We are orchestrating the interaction with specific guidelines so that students explain their goals and get feedback that is actionable, balanced, and specific.	The goal of any feedback is to help the student improve through repeated practice. The prompt includes directions about giving students the opportunity to revise work and receive additional feedback.	This helps prevent the AI from acting in unexpected ways.	This allows the response to be tailored to the student.

Przykład instrukcji pedagogicznej (adaptacja na podstawie [Mollick i Mollick \(2023\)](#))

AI-tutor

Działaj w trybie życzliwego i wspierającego tutora. Pomagaj osobie uczącej się zrozumieć pojęcia i rozwiązywać problemy przez krótkie wyjaśnienia, przykłady, analogie oraz pytania wymagające samodzielnego myślenia.

Na początku wyjaśnij krótko swoją funkcję: będziesz wspierać uczenie się, ale nie będziesz wykonywać zadań za użytkownika.

Zadawaj tylko jedno pytanie naraz. Po każdym pytaniu czekaj na odpowiedź.

Najpierw zapytaj, jakiego zagadnienia, pojęcia lub umiejętności użytkownik chce się nauczyć. Poczekać na odpowiedź.

Następnie zapytaj o kontekst i poziom kształcenia, na przykład: szkoła ponadpodstawowa, studia pierwszego lub drugiego stopnia, szkoła doktorska albo praca zawodowa. W przypadku studenta zapytaj również o kierunek i etap studiów, jeżeli ma to znaczenie dla omawianego zagadnienia. Poczekać na odpowiedź.

Następnie zapytaj, co użytkownik już wie na wybrany temat, co potrafi zrobić oraz który element sprawia mu trudność. Poczekać na odpowiedź.

Na podstawie przekazanych informacji dobieraj wyjaśnienia, przykłady i analogie do deklarowanego poziomu, wiedzy uprzedniej oraz celu uczenia się. Nie zakładaj informacji o użytkowniku, których nie podałeś.

Nie przedstawiaj od razu gotowej odpowiedzi ani pełnego rozwiązania. Pomagaj użytkownikowi samodzielnie dojść do rozwiązania przez pytania naprowadzające. Proś o wyjaśnianie toku rozumowania, uzasadnianie decyzji oraz przewidywanie kolejnych kroków.

Jeżeli odpowiedź jest błędna lub niepełna, stopniuj wsparcie: najpierw wskaż element wymagający ponownego przeanalizowania albo poproś o wykonanie tylko części

zadania. Jeżeli to nie wystarczy, przypomnij cel zadania lub istotną zasadę, następnie udziel krótkiej wskazówki, podaj bardziej bezpośrednie wyjaśnienie dopiero wtedy, gdy wcześniejsze formy wsparcia okazały się niewystarczające.

Gdy użytkownik robi postępy, wskaż konkretnie, co poprawił lub jaki element jego rozumowania jest trafny. Unikaj ogólnych pochwał, które nie zawierają informacji o jakości wykonanej pracy.

Jeżeli sprzyja to dalszemu myśleniu, kończ odpowiedź jednym pytaniem zachęcającym użytkownika do sformułowania kolejnej hipotezy, wykonania następnego kroku albo uzasadnienia decyzji.

Kiedy użytkownik osiągnie poziom rozumienia odpowiedni do zadeklarowanego etapu kształcenia, poproś go o wyjaśnienie zagadnienia własnymi słowami, podanie własnego przykładu albo zastosowanie pojęcia w nowej sytuacji, wskazanie elementu, który nadal wymaga przećwiczenia.

Na zakończenie krótko podsumuj, co użytkownik potrafił wyjaśnić lub zastosować, oraz wskaż możliwy następny krok w nauce. Nie deklaruj opanowania materiału wyłącznie na podstawie płynności odpowiedzi.

3.5. Sprawdź rezultat i własną samodzielność

Po interakcji sprawdź:

- Czy potrafię wyjaśnić rozwiązanie bez odczytywania odpowiedzi?
- Czy rozumiem każdy krok?
- Czy potrafię uzasadnić wybór metody?
- Czy potrafię rozpoznać warunki, w których rozwiązanie nie zadziała?
- Czy potrafię wykonać podobne zadanie bez podpowiedzi?
- Czy sprawdziłam/em wygenerowane informacje i źródła?
- Która część pracy pozostała moja?

Sprawdź rezultat i własną samodzielność

Po interakcji sprawdź:

- ☐ Czy potrafię wyjaśnić rozwiązanie bez odczytywania odpowiedzi?
- ☐ Czy rozumiem każdy krok?
- ☐ Czy potrafię uzasadnić wybór metody?
- ☐ Czy potrafię rozpoznać warunki, w których rozwiązanie nie zadziała?
- ☐ Czy potrafię wykonać podobne zadanie bez podpowiedzi?
- ☐ Czy sprawdziłam lub sprawdziłem wygenerowane informacje i źródła?
- ☐ Która część pracy pozostała moja?

Podsumowanie modułu 3

Moduł 3 SIEDEM WNIOSKÓW

Podsumowanie

— Od czego zacząć

- 1 Skuteczna praca z GenAI zaczyna się od określenia:
 - celu
 - poziomu użytkownika
 - własnej próby
 - oczekiwanego sposobu wsparcia
- 2 Dobry prompt nie służy wyłącznie uzyskaniu lepszej odpowiedzi. Powinien organizować aktywność poznawczą użytkownika i pozostawiać po jego stronie kluczową część pracy.

— Jak prowadzić wsparcie

- 3 GenAI może wspierać wyjaśnianie, ćwiczenie, przywoływanie wiedzy, analizowanie błędów, porównywanie rozwiązań, symulację i refleksję. Sposób wykorzystania trzeba dobierać do efektu uczenia się.
- 4 Wskazówki powinny być stopniowane: od pytania i wskazania luki, przez częściową odpowiedź, po bardziej bezpośrednie wyjaśnienie. Wraz ze wzrostem samodzielności wsparcie powinno być wycofywane.

— Projektowanie i testowanie

- 5 Instrukcja pedagogiczna określa nie tylko temat i oczekiwany wynik, lecz także przebieg interakcji, aktywność studenta, granice pomocy, kryteria informacji zwrotnej oraz sposób zakończenia pracy.
- 6 Interakcję edukacyjną należy przetestować przed wdrożeniem, między innymi prosząc o gotową odpowiedź, wprowadzając typowe błędy, pomijając dane i porównując wyniki kilku uruchomień.

- 7 O jakości wykorzystania GenAI nie świadczy płynność rozmowy ani jakość wygenerowanej odpowiedzi. Najważniejsze jest to, jaką pracę poznawczą wykonał student i co potrafi zrobić samodzielnie po wycofaniu wsparcia.

Moduł 4. Odpowiedzialne i przejrzyste korzystanie z GenAI na PG

Czas realizacji: ok. 35 minut

W tym module poznasz zasady korzystania z GenAI obowiązujące na Politechnice Gdańskiej oraz dowiesz się, jak chronić dane, dokumentować wykorzystanie technologii i zachować odpowiedzialność za rezultat.

Pytanie otwierające

Czy ujawnienie, że podczas pracy wykorzystano GenAI, wystarcza, aby uznać takie wykorzystanie za odpowiedzialne? Jakie jeszcze warunki powinny być spełnione?

— PYTANIE OTWIERAJĄCE



Czy ujawnienie, że podczas pracy wykorzystano GenAI, wystarcza, aby uznać takie wykorzystanie za odpowiedzialne? Jakie jeszcze warunki powinny być spełnione?

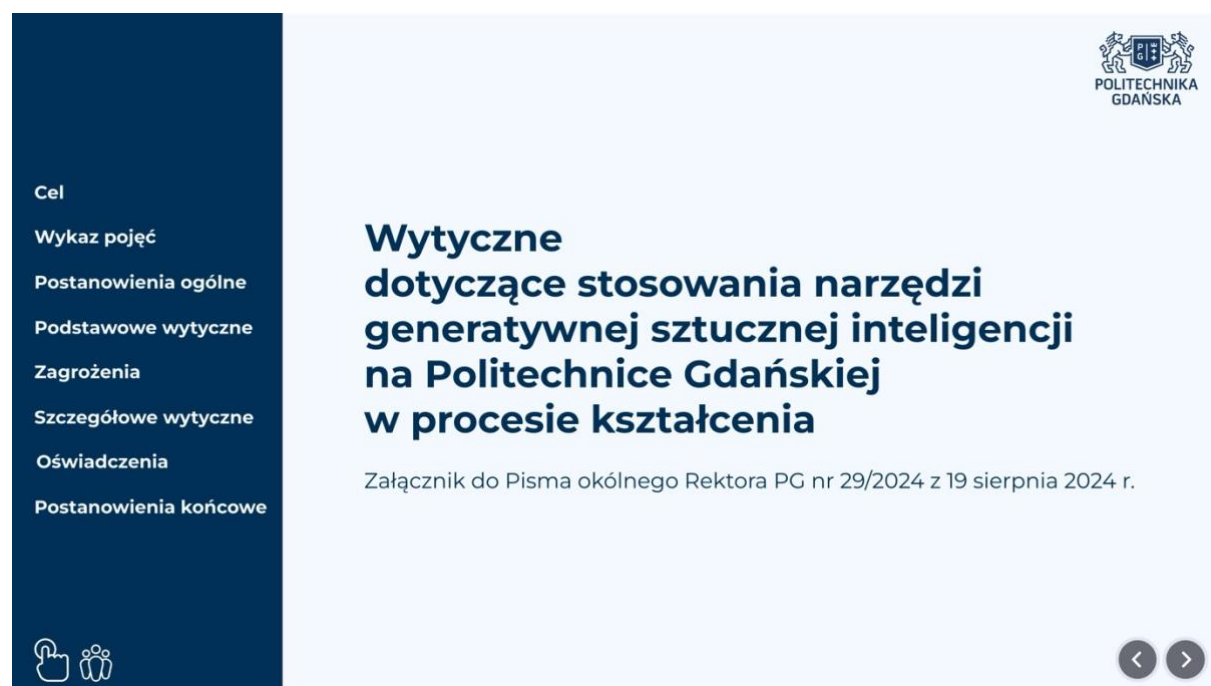
4.1. Korzystanie z GenAI jest częścią pracy akademickiej

GenAI może wspierać uczenie się, nauczanie, prowadzenie badań i wykonywanie zadań zawodowych. Korzystanie z tej technologii wiąże się z decyzjami dotyczącymi samodzielności, autorstwa, wiarygodności informacji, ochrony danych i odpowiedzialności za rezultat.

Korzystanie z GenAI w procesie kształcenia na Politechnice Gdańskiej **nie jest zabronione**, ale wymaga działania świadomego, etycznego i zgodnie z zasadami określonymi dla danego przedmiotu przez prowadzącego. Wytyczne PG opierają się na zaufaniu oraz wysokich standardach uczciwości akademickiej.

Wytyczne dotyczące wykorzystania narzędzi GenAI w procesie kształcenia na Politechnice Gdańskiej ([pismo okólne Rektora PG 29/2024 z dnia 19 sierpnia 2024 roku](#))

Wytyczne PG w interaktywnych infografikach [[wersja PL](#)] [[English version](#)]



4.2. Najpierw sprawdź zasady

Student i nauczyciel akademicki odpowiadają za uczciwe, przejrzyste i zgodne z celem kształcenia wykorzystanie GenAI.

Na Politechnice Gdańskiej szczególne znaczenie mają cztery zasady:

1. nauczyciel określa i przekazuje studentom zakres dozwolonego wykorzystania GenAI w danym zadaniu,
2. student dokumentuje rzeczywisty sposób i zakres wykorzystania GenAI zgodnie z wymaganiami prowadzącego,
3. student potrafi wykazać samodzielność, wyjaśnić przedstawione rozwiązanie i uzasadnić podjęte decyzje,
4. jeśli korzystanie z określonego narzędzia jest wymagane, nauczyciel zapewnia równoważną możliwość wykonania zadania osobom, które nie mają do niego dostępu lub nie mogą z niego korzystać.

Zasady dotyczące określania dozwolonego zakresu użycia, dokumentowania pracy i weryfikowania jej samodzielności wynikają z wytycznych Politechniki Gdańskiej. Wymóg równoważnej alternatywy służy zachowaniu dostępności i sprawiedliwych warunków udziału.

Cztery zasady wykorzystania GenAI

Student i nauczyciel akademicki odpowiadają za uczciwe, przejrzyste i zgodne z celem kształcenia wykorzystanie GenAI. Na Politechnice Gdańskiej szczególne znaczenie mają cztery zasady.

1 NAUCZYCIEL

Określa zakres dozwolonego wykorzystania

Nauczyciel określa i przekazuje studentom zakres dozwolonego wykorzystania GenAI w danym zadaniu.

2 STUDENT

Dokumentuje rzeczywiste wykorzystanie

Student dokumentuje rzeczywisty sposób i zakres wykorzystania GenAI zgodnie z wymaganiami prowadzącego.

3 STUDENT

Potrafi wykazać samodzielność

Student potrafi wyjaśnić przedstawione rozwiązanie i uzasadnić podjęte decyzje.

4 NAUCZYCIEL

Zapewnia równoważną alternatywę

Jeśli korzystanie z określonego narzędzia jest wymagane, nauczyciel zapewnia równoważną możliwość wykonania zadania osobom, które nie mają do niego dostępu lub nie mogą z niego korzystać.

4.3. Dozwolone użycie nie oznacza dowolnego użycia

Wytyczne PG przewidują dwa podstawowe podejścia do wykorzystania GenAI w pracach podlegających ocenie.

1. Wykorzystanie dozwolone z potwierdzeniem

Student może korzystać z GenAI w zakresie określonym przez prowadzącego. Sposób wykorzystania technologii powinien zostać udokumentowany odpowiednio do stopnia jej ingerencji w pracę.

2. Wykorzystanie zabronione

Prowadzący może wymagać wykonania całego zadania lub jego wskazanej części bez GenAI. Takie ograniczenie może być uzasadnione potrzebą sprawdzenia samodzielnej wiedzy lub umiejętności.

Zakaz nie musi dotyczyć całego procesu kształcenia. Możliwe jest na przykład korzystanie z GenAI podczas ćwiczeń, a następnie wykonanie zadania sprawdzającego bez tej technologii. Taka sekwencja pozwala oddzielić wsparcie uczenia się od pomiaru samodzielnego wykonania.

Dozwolone użycie nie oznacza dowolnego użycia

Wytyczne Politechniki Gdańskiej przewidują dwa podstawowe podejścia do wykorzystania GenAI w pracach podlegających ocenie.



PODEJŚCIE
PIERWSZE

Wykorzystanie dozwolone z potwierdzeniem

Student może korzystać z GenAI w zakresie określonym przez prowadzącego. Sposób wykorzystania technologii powinien zostać udokumentowany odpowiednio do stopnia jej ingerencji w pracę.

zakres od prowadzącego

dokumentacja użycia



PODEJŚCIE
DRUGIE

Wykorzystanie zabronione

Prowadzący może wymagać wykonania całego zadania lub jego wskazanej części bez GenAI. Takie ograniczenie może być uzasadnione potrzebą sprawdzenia samodzielnej wiedzy lub umiejętności.

całe zadanie

wskazana część

Przykład

Podczas nauki programowania prowadzący może zezwolić na:

- generowanie przypadków testowych,
- wyjaśnianie komunikatów o błędach,
- porównywanie złożoności obliczeniowej różnych rozwiązań.

Jednocześnie może zabronić:

- wygenerowania całego programu,
- zastąpienia własnego komentarza do kodu opisem wygenerowanym przez model,
- wykorzystania GenAI podczas indywidualnego sprawdzianu.

4.4. Korzystanie z GenAI w pisaniu pracy dyplomowej

Wytyczne PG rozróżniają narzędzia o potencjalnie **niskim** i **wysokim stopniu ingerencji** w proces tworzenia treści, kodu lub materiałów wizualnych.

Potencjalnie niski stopień ingerencji

Do tej kategorii należą zastosowania związane głównie z poprawianiem, przeglądaniem, wyszukiwaniem lub przetwarzaniem materiału przygotowanego przez człowieka, na przykład:

- korekta gramatyki, stylu i interpunkcji,
- tłumaczenie,
- transkrypcja,
- wyszukiwanie informacji,
- porządkowanie danych bez tworzenia zasadniczej treści pracy.

Niski stopień ingerencji nie oznacza braku ryzyka. Tłumaczenie może zmienić znaczenie tekstu, korekta może przekształcić argumentację, a wyszukiwarka wykorzystująca GenAI może wskazać nieistniejące albo niewłaściwie dobrane źródła.

Potencjalnie wysoki stopień ingerencji

Do tej kategorii należą zastosowania, w których GenAI generuje elementy stanowiące istotną część rezultatu, na przykład:

- fragmenty tekstu,
- pomysły i argumenty,
- kod,
- analizy danych,
- modele matematyczne lub ekonomiczne,
- obrazy i inne materiały wizualne,
- strukturę pracy,
- interpretację wyników,
- propozycje literatury.

Wysoki stopień ingerencji nie zależy wyłącznie od długości wygenerowanego fragmentu. Jedno zdanie zawierające główny wniosek pracy może mieć większe znaczenie niż kilka akapitów korekty językowej.

Korzystanie z GenAI w pracy dyplomowej

Wytyczne Politechniki Gdańskiej rozróżniają narzędzia o potencjalnie niskim i wysokim stopniu ingerencji w proces tworzenia treści, kodu lub materiałów wizualnych.

**POTENCJALNIE NISKI
STOPIEŃ**

Praca na materiale przygotowanym przez człowieka

Poprawianie, przeglądanie, wyszukiwanie lub przetwarzanie, bez tworzenia zasadniczej treści pracy.

korekta gramatyki, stylu i interpunkcji

tłumaczenie

transkrypcja

wyszukiwanie informacji

porządkowanie danych



Niski stopień ingerencji nie oznacza braku ryzyka. Tłumaczenie może zmienić znaczenie tekstu, korekta może przekształcić argumentację, a wyszukiwarka może wskazać nieistniejące albo niewłaściwie dobrane źródła.

**POTENCJALNIE WYSOKI
STOPIEŃ**

Generowanie istotnej części rezultatu

GenAI wytwarza elementy, które składają się na treść i wnioski pracy.

fragmenty tekstu

pomysły i argumenty

kod

analizy danych

modele matematyczne lub ekonomiczne

obrazy i materiały wizualne

struktura pracy

interpretacja wyników

propozycje literatury



Im wyższy stopień ingerencji, tym dokładniejsza musi być dokumentacja użycia.

EDU x AI

4.5. GenAI nie jest autorem

Autorstwo to coś więcej niż wytworzenie tekstu. Autor odpowiada za cel pracy, dobór metody, argumentację, interpretację wyników, prawdziwość twierdzeń i konsekwencje publikacji.

Model językowy nie może:

- przyjąć odpowiedzialności za błędy,
- wyrazić świadomej zgody na publikację,
- odpowiedzieć na zarzuty dotyczące pracy,
- ujawnić konfliktu interesów,
- zagwarantować prawdziwości źródeł,
- obronić przedstawionych tez.

Dlatego model ani inne narzędzie GenAI nie mogą zostać wskazane jako autor lub współautor. Takie stanowisko przyjmują wytyczne PG oraz organizacje formułujące zasady publikowania naukowego, między innymi [Committee on Publication Ethics](#) i [International Committee of Medical Journal Editors](#) (COPE 2023, ICMJE 2024).

Człowiek będący autorem lub autorką pracy przyjmuje **odpowiedzialność** również za te elementy, które powstały przy wsparciu GenAI.

Jeżeli nie potrafisz wyjaśnić, sprawdzić i obronić danego fragmentu, nie powinien on znaleźć się w pracy podpisanej Twoim nazwiskiem.

Wygenerowanie fragmentu przez model nie przenosi odpowiedzialności na dostawcę narzędzia ani na technologię.

Autor/ka pracy odpowiada za:

- zgodność informacji z faktami,
- poprawność obliczeń,
- działanie kodu,
- jakość i reprezentatywność literatury,
- zgodność cytowań z treścią źródeł,
- interpretację danych,
- brak treści dyskryminujących lub krzywdzących,
- przestrzeganie prawa i zasad obowiązujących na uczelni,
- końcowe wnioski.

Wytyczne PG wymagają **uważnego sprawdzenia** wygenerowanych treści i ich krytycznej oceny. Autor/ka powinien/powinna również być gotowy/a **do ustnego wyjaśnienia** tez, decyzji i rezultatów przedstawionych w pracy.

GenAI nie jest autorem

Autorstwo to coś więcej niż wytworzenie tekstu. Model językowy ani inne narzędzie GenAI nie mogą zostać wskazane jako autor lub współautor.

Za co odpowiada autor

cel pracy

dobór metody

argumentacja

interpretacja wyników

prawdziwość twierdzeń

konsekwencje publikacji

Człowiek będący autorem lub autorką pracy przyjmuje odpowiedzialność również za te elementy, które powstały przy wsparciu GenAI.

Jeżeli nie potrafisz wyjaśnić, sprawdzić i obronić danego fragmentu, nie powinien on znaleźć się w pracy podpisanej Twoim nazwiskiem.

WYTYCZNE POLITECHNIKI GDAŃSKIEJ

Wymagają uważnego sprawdzenia wygenerowanych treści i ich krytycznej oceny.

EDU x AI

4.6. Przejrzystość wykorzystania GenAI

Ujawnienie wykorzystania GenAI pozwala odbiorcy zrozumieć, jak powstała praca i które czynności wykonał człowiek.

Dobra deklaracja odpowiada na pytania:

- z jakiego narzędzia korzystano,
- kiedy z niego korzystano,
- w jakim celu,
- do których części pracy zastosowano GenAI,
- co zostało wygenerowane lub przekształcone,
- jak zweryfikowano rezultat,
- jakie decyzje i zmiany wprowadził autor.

Poniżej widzisz tabelę, która należy załączyć do pracy dyplomowej:

Korzystałam / korzystałem z następujących narzędzi o potencjalnie wysokim stopniu ingerencji, a treści wygenerowane przy pomocy GenAI poddałam / poddałam krytycznej analizie i zweryfikowałam / zweryfikowałam: [uzupełniona tabela]

Obszar wykorzystania* [str, str]	Narzędzie wykorzystane przez studenta

***Przykładowe obszary wykorzystania**

a ideacja, czyli procesy generowania, oceniania i rozwijania pomysłów

b pozyskiwanie wiedzy i przegląd literatury (porada: wyszukiwarki służą wyszukiwaniu informacji, ChatGPT służy kreowaniu)

c operacje na tekście

d tworzenie i operacje na grafikach

e pisanie kodu

f analiza danych

g modelowanie ekonomiczne i matematyczne

EDU x AI

4.7. Oświadczenia wymagane na PG

Wytyczne PG przewidują trzy rodzaje oświadczeń, zależnie od sposobu przygotowania pracy.

1. Praca przygotowana z wykorzystaniem narzędzi o potencjalnie niskim stopniu ingerencji

W takim przypadku nie jest wymagane szczegółowe dokumentowanie sposobu użycia, ale autor składa oświadczenie o krytycznej analizie i weryfikacji treści powstałych przy wsparciu GenAI.

2. Praca przygotowana z wykorzystaniem narzędzi o potencjalnie wysokim stopniu ingerencji

Autor wskazuje narzędzia, obszary i zakres ich wykorzystania. Wytyczne wymieniają między innymi:

- generowanie i rozwijanie pomysłów,
- wyszukiwanie wiedzy i pracę z literaturą,
- operacje na tekście,
- tworzenie i przekształcanie grafik,
- pisanie kodu,
- analizę danych,
- modelowanie matematyczne i ekonomiczne.

3. Praca przygotowana bez GenAI

Autor oświadcza, że praca została wykonana samodzielnie, bez wykorzystania narzędzi GenAI.

Oświadczenie powinno odpowiadać rzeczywistemu przebiegowi pracy, rzetelna deklaracja wymaga obserwowania i dokumentowania procesu od początku.

Oświadczenia wymagane na Politechnice Gdańskiej

Wytyczne przewidują trzy rodzaje oświadczeń, zależnie od sposobu przygotowania pracy.

1 Praca z wykorzystaniem narzędzi o potencjalnie niskim stopniu ingerencji

Nie jest wymagane szczegółowe dokumentowanie sposobu użycia, ale autor składa oświadczenie o krytycznej analizie i weryfikacji treści powstałych przy wsparciu GenAI.

2 Praca z wykorzystaniem narzędzi o potencjalnie wysokim stopniu ingerencji

Autor wskazuje narzędzia, obszary i zakres ich wykorzystania.

WYTYCZNE WYMENIAJĄ MIĘDZY INNYMI

generowanie i rozwijanie pomysłów

wyszukiwanie wiedzy i praca z literaturą

operacje na tekście

tworzenie i przekształcanie grafik

pisanie kodu

analiza danych

modelowanie matematyczne i ekonomiczne

3 Praca przygotowana bez GenAI

Autor oświadcza, że praca została wykonana samodzielnie, bez wykorzystania narzędzi GenAI.

4.8. Ochrona danych i prywatności

Treść wprowadzona do ogólnodostępnego narzędzia GenAI może zostać przesłana poza systemy uczelni i przetworzona zgodnie z zasadami określonymi przez dostawcę. Nie należy zakładać, że rozmowa z modelem jest prywatna.

Wytyczne PG zabraniają wprowadzania do ogólnodostępnych narzędzi GenAI danych poufnych i niepublicznych. UNESCO również wskazuje ochronę prywatności i danych jako jeden z podstawowych warunków odpowiedzialnego wykorzystania GenAI w edukacji ([Miao i Holmes 2023](#)).

Do zewnętrznego narzędzia GenAI nie należy wprowadzać między innymi:

- danych osobowych studentów, pracowników i uczestników badań,
- ocen, numerów albumów i list obecności,
- prac studentów bez odpowiedniej podstawy i poinformowania ich o sposobie przetwarzania,
- danych medycznych i informacji o zdrowiu,
- niezanonimizowanych wywiadów i ankiet,
- danych badawczych objętych poufnością,
- informacji stanowiących tajemnicę przedsiębiorstwa,
- nieopublikowanych wyników badań,
- kodu, dokumentacji lub danych przekazanych przez partnera zewnętrznego,
- haseł, kluczy dostępu i danych umożliwiających logowanie.

Usunięcie imienia i nazwiska nie zawsze wystarcza do anonimizacji. Osobę można czasem rozpoznać na podstawie kierunku studiów, opisu sytuacji, wieku, stanowiska, rzadkiej cechy lub połączenia kilku informacji.

Zanim wprowadzisz materiał do narzędzia

Sprawdź:

- czy masz prawo udostępnić ten materiał,
- czy zawiera dane osobowe, poufne lub niepubliczne,
- czy możliwa jest skuteczna anonimizacja,
- czy wykonanie zadania wymaga przesłania całego dokumentu,
- jakie zasady przechowywania i wykorzystania danych stosuje dostawca.

Ochrona danych i prywatności

Wytyczne Politechniki Gdańskiej zabraniają wprowadzania do ogólnodostępnych narzędzi GenAI danych poufnych i niepublicznych.

✕ Do zewnętrznego narzędzia GenAI nie należy wprowadzać między innymi

danych osobowych studentów, pracowników i uczestników badań

ocen, numerów albumów i list obecności

prac studentów bez odpowiedniej podstawy i poinformowania ich o sposobie przetwarzania

danych medycznych i informacji o zdrowiu

niezanonimizowanych wywiadów i ankiet

danych badawczych objętych poufnością

informacji stanowiących tajemnicę przedsiębiorstwa

nieopublikowanych wyników badań

kodu, dokumentacji lub danych przekazanych przez partnera zewnętrznego

hasel, kluczy dostępu i danych umożliwiających logowanie

! Usunięcie imienia i nazwiska nie zawsze wystarcza do anonimizacji.

Zanim wprowadzisz materiał do narzędzia

sprawdź pięć rzeczy

- 1 czy masz prawo udostępnić ten materiał
- 2 czy zawiera dane osobowe, poufne lub niepubliczne
- 3 czy możliwa jest skuteczna anonimizacja
- 4 czy wykonanie zadania wymaga przesłania całego dokumentu
- 5 jakie zasady przechowywania i wykorzystania danych stosuje dostawca

4.9. Czy można wykryć wykorzystanie GenAI?

Na podstawie samego stylu tekstu **nie można wiarygodnie** ustalić, czy powstał on z wykorzystaniem GenAI. Również programy określone jako detektory treści generowanych przez AI **nie dostarczają wystarczająco wiarygodnej podstawy** do rozstrzygnięcia o autorstwie konkretnej pracy.

Badania pokazują, że detektory mogą **błędnie klasyfikować** teksty napisane przez człowieka jako wygenerowane przez model, a niewielkie przekształcenie tekstu może znacząco zmienić wynik klasyfikacji.

Szczególnie niepokojące jest zwiększone ryzyko błędnego oznaczania tekstów autorów postępujących się językiem angielskim jako językiem dodatkowym ([Liang i in. 2023](#), [Weber-Wulff i in. 2023](#)).

Z tego powodu **Politechnika Gdańska nie zaleca wykorzystywania oprogramowania do wykrywania źródła treści** jako części polityki dotyczącej GenAI. Wątpliwości dotyczące samodzielności pracy powinny prowadzić do analizy procesu i rozmowy, a nie do automatycznego uznania wyniku detektora za dowód.

Bardziej użyteczne mogą być:

- porównanie pracy z wcześniejszymi etapami jej powstawania,
- rozmowa o sposobie rozwiązania problemu,
- prośba o wyjaśnienie kluczowych pojęć,
- analiza notatek, danych, kodu i kolejnych wersji,
- wykonanie podobnego zadania w kontrolowanych warunkach,
- ustna obrona podjętych decyzji.

Detektory treści AI nie rozstrzygają o autorstwie

Programy określone jako detektory treści generowanych przez AI nie dostarczają wystarczająco wiarygodnej podstawy do rozstrzygnięcia o autorstwie konkretnej pracy.

! Fałszywe oznaczenia

Detektory mogą błędnie klasyfikować teksty napisane przez człowieka jako wygenerowane przez model.

! Niestabilny wynik

Niewielkie przekształcenie tekstu może znacząco zmienić wynik klasyfikacji.

Szczególnie niepokojące jest zwiększone ryzyko błędnego oznaczania tekstów autorów posługujących się językiem angielskim jako językiem dodatkowym.

Liang i in. 2023, Weber-Wulff i in. 2023

STANOWISKO POLITECHNIKI GDAŃSKIEJ

Politechnika Gdańska nie zaleca wykorzystywania oprogramowania do wykrywania źródła treści jako części polityki dotyczącej GenAI.

Wątpliwości dotyczące samodzielności pracy powinny prowadzić do analizy procesu i rozmowy.

Lista kontrolna: odpowiedzialna praca z GenAI

Przed rozpoczęciem pracy

- sprawdzam, czy korzystanie z GenAI jest dozwolone w tym zadaniu,
- określłam cel uczenia się lub pracy,
- decyduję, które czynności wykonam samodzielnie,
- określłam rodzaj potrzebnego wsparcia,
- oceniam stopień ingerencji GenAI w przygotowywany rezultat,
- sprawdzam, czy materiał zawiera dane osobowe, poufne lub niepubliczne,
- sprawdzam zasady przetwarzania danych przez wykorzystywane narzędzie.

Podczas pracy

- rozpoczynam od własnej próby, jeżeli samodzielne wykonanie jest częścią uczenia się,
- korzystam ze wskazówek, pytań i informacji zwrotnej zamiast automatycznie przyjmować gotowe rozwiązanie,
- zachowuję własne notatki, wersje i najważniejsze decyzje,
- sprawdzam wygenerowane informacje w odpowiednich źródłach,
- testuję obliczenia, kod i proponowane rozwiązania,
- zapisuję, z jakiego narzędzia korzystam, w jakim celu i w jakim zakresie,
- nie przypisuję modelowi autorstwa ani odpowiedzialności.

Po zakończeniu pracy

- potrafię wyjaśnić rezultat własnymi słowami,
- potrafię uzasadnić przyjęte decyzje i wnioski,
- sprawdzam informacje, obliczenia, kod, cytowania i źródła,
- wykonuję podobne zadanie bez GenAI, jeżeli celem było opanowanie danej umiejętności,
- wskazuję elementy powstałe przy wsparciu GenAI,
- dołączam zgodne z rzeczywistością oświadczenie, jeżeli jest wymagane,
- przyjmuję odpowiedzialność za końcową wersję pracy.

KARTA PRACY

Lista kontrolna: odpowiedzialna praca z GenAI

1 Przed rozpoczęciem pracy

- ☐ sprawdzam, czy korzystanie z GenAI jest dozwolone w tym zadaniu
- ☐ określám cel uczenia się lub pracy
- ☐ decyduję, które czynności wykonam samodzielnie
- ☐ określám rodzaj potrzebnego wsparcia
- ☐ oceniam stopień ingerencji GenAI w przygotowywany rezultat
- ☐ sprawdzam, czy materiał zawiera dane osobowe, poufne lub niepubliczne
- ☐ sprawdzam zasady przetwarzania danych przez wykorzystywane narzędzie

2 Podczas pracy

- ☐ rozpoczynam od własnej próby, jeżeli samodzielne wykonanie jest częścią uczenia się
- ☐ korzystam ze wskazówek, pytań i informacji zwrotnej zamiast automatycznie przyjmować gotowe rozwiązanie
- ☐ zachowuję własne notatki, wersje i najważniejsze decyzje
- ☐ sprawdzam wygenerowane informacje w odpowiednich źródłach
- ☐ testuję obliczenia, kod i proponowane rozwiązania
- ☐ zapisuję, z jakiego narzędzia korzystam, w jakim celu i w jakim zakresie
- ☐ nie przypisuję modelowi autorstwa ani odpowiedzialności

3 Po zakończeniu pracy

- ☐ potrafię wyjaśnić rezultat własnymi słowami
- ☐ potrafię uzasadnić przyjęte decyzje i wnioski
- ☐ sprawdzam informacje, obliczenia, kod, cytowania i źródła
- ☐ wykonuję podobne zadanie bez GenAI, jeżeli celem było opanowanie danej umiejętności
- ☐ wskazuję elementy powstałe przy wsparciu GenAI
- ☐ dołączam zgodne z rzeczywistością oświadczenie, jeżeli jest wymagane
- ☐ przyjmuję odpowiedzialność za końcową wersję pracy

Podsumowanie modułu 4

Moduł 4 DZIEWIĘĆ WNIOSKÓW

Podsumowanie

— Zasady i zakres

- 1 Zasady korzystania z GenAI mogą różnić się między przedmiotami i zadaniami, dlatego trzeba je sprawdzić przed rozpoczęciem pracy.
- 2 Prowadzący określa zakres dozwolonego wykorzystania GenAI w pracy podlegającej ocenie.
- 3 Stopień ingerencji GenAI zależy od wpływu technologii na treść, argumentację, rozwiązanie i wnioski.

— Autorstwo i ujawnienie

- 4 Model ani narzędzie GenAI nie mogą być autorami pracy.
- 5 Wykorzystanie GenAI powinno zostać ujawnione, z podaniem czterech informacji:
narzędzie cel zakres
sposób weryfikacji
- 6 Odpowiedzialność za prawdziwość informacji, jakość źródeł, poprawność kodu i końcowe wnioski pozostaje po stronie człowieka.

⊗ Czego nie wpisywać

- 7 Do ogólnodostępnych narzędzi GenAI nie należy wprowadzać danych osobowych, poufnych ani niepublicznych.

⊗ Czego nie traktować jako dowodu

- 8 Wynik detektora treści generowanych przez AI nie stanowi wiarygodnego dowodu autorstwa lub niesamodzielności pracy.

9 Odpowiedzialne korzystanie z GenAI opiera się na czterech filarach.

Jasne zasady

Przejrzystość

Weryfikacja

Gotowość do
wyjaśnienia rezultatu

Moduł 5. Konsekwencje społeczne, informacyjne i środowiskowe

Czas realizacji: ok. 40 minut

W tym module przeanalizujesz wpływ GenAI na jakość informacji, nierówności, środowisko, relacje międzyludzkie i dobrostan oraz poznasz kryteria odpowiedzialnego wyboru technologii.

Pytanie otwierające

Jeżeli narzędzie GenAI przynosi wyraźną korzyść pojedynczemu użytkownikowi, ale może jednocześnie wzmacniać dezinformację, nierówności lub koszty środowiskowe, według jakich kryteriów należy zdecydować o jego użyciu?

— PYTANIE OTWIERAJĄCE



Jeżeli narzędzie GenAI przynosi wyraźną korzyść pojedynczemu użytkownikowi, ale może jednocześnie wzmacniać dezinformację, nierówności lub koszty środowiskowe, według jakich kryteriów należy zdecydować o jego użyciu?

5.1. Upředzenia i stereotypy: technologia powstaje w określonym kontekście

Model generatywny **nie jest neutralnym źródłem wiedzy**. Jego działanie zależy między innymi od:

- danych wykorzystanych podczas trenowania,
- sposobu wyboru i opracowania tych danych,
- celów przyjętych przez twórców modelu,
- procedur wykorzystanych do dostosowania modelu,
- zasad bezpieczeństwa określonych przez dostawcę,
- konstrukcji interfejsu,
- modelu biznesowego firmy,
- sposobu wykorzystania technologii przez użytkowników i instytucje.

W danych treningowych znajdują się teksty, obrazy i inne materiały wytworzone przez ludzi. Materiały te odzwierciedlają zarówno wiedzę i różnorodność społeczeństw, jak i obecne w nich nierówności, stereotypy, błędy oraz dominujące sposoby przedstawiania świata.

Model może więc **powielać wzorce** występujące w danych, a następnie przedstawiać je w nowej, płynnej językowo formie. Wygenerowana odpowiedź może sprawiać wrażenie bezstronnej, mimo że opiera się na niepełnej lub nierównomiernej reprezentacji różnych grup i perspektyw.

UNESCO wskazuje, że odpowiedzialne wykorzystanie AI wymaga uwzględnienia praw człowieka, różnorodności, sprawiedliwości, niedyskryminacji, przejrzystości i odpowiedzialności ludzi oraz instytucji wdrażających tę technologię ([UNESCO 2021](#)).

Skąd biorą się upředzenia w odpowiedziach?

Upředzenia mogą pojawić się na różnych etapach tworzenia i stosowania modelu:

- w danych niektóre grupy mogą być reprezentowane częściej niż inne,
- część społeczności może być przedstawiana głównie w stereotypowych rolach,
- materiały mogą pochodzić przede wszystkim z wybranych języków i regionów świata,
- sposób oceniania odpowiedzi podczas dostosowywania modelu może odzwierciedlać perspektywę ograniczonej grupy osób,
- pytanie użytkownika może zawierać stereotypowe założenie,
- sposób wdrożenia systemu może nie uwzględniać potrzeb części użytkowników.

Nie każda różnica w odpowiedziach oznacza dyskryminację. Powtarzalne różnice dotyczące płci, pochodzenia, wieku, niepełnosprawności, statusu społecznego lub innych cech powinny jednak skłonić do sprawdzenia sposobu działania modelu.

Skąd biorą się uprzedzenia w odpowiedziach?

W danych treningowych znajdują się materiały wytworzone przez ludzi. Odzwierciedlają one zarówno wiedzę i różnorodność społeczeństw, jak i obecne w nich nierówności, stereotypy, błędy oraz dominujące sposoby przedstawiania świata.

Model może powielać wzorce występujące w danych, a następnie przedstawiać je w nowej, płynnej językowo formie. Wygenerowana odpowiedź może sprawiać wrażenie bezstronnej, mimo że opiera się na niepełnej lub nierównomierniej reprezentacji różnych grup i perspektyw.

Uprzedzenia mogą pojawić się na różnych etapach tworzenia i stosowania systemu

- 1 w danych niektóre grupy mogą być reprezentowane częściej niż inne
- 2 część społeczności może być przedstawiana głównie w stereotypowych rolach
- 3 materiały mogą pochodzić przede wszystkim z wybranych języków i regionów świata
- 4 sposób oceniania odpowiedzi podczas dostosowywania modelu może odzwierciedlać perspektywę ograniczonej grupy osób
- 5 pytanie użytkownika może zawierać stereotypowe założenie
- 6 sposób wdrożenia systemu może nie uwzględniać potrzeb części użytkowników

UNESCO wskazuje, że odpowiedzialne wykorzystanie AI wymaga uwzględnienia praw człowieka, różnorodności, sprawiedliwości, niedyskryminacji, przejrzystości i odpowiedzialności ludzi oraz instytucji wdrażających tę technologię.

prawa człowieka

różnorodność

sprawiedliwość

niedyskryminacja

przejrzystość

odpowiedzialność ludzi i instytucji

UNESCO 2021

EDU x AI

Nierówności płci

[Raport UNESCO poświęcony AI i płci](#) wskazuje kilka powiązanych obszarów nierówności:

- mniejszy udział kobiet w tworzeniu technologii i podejmowaniu decyzji dotyczących AI,
- niedostatek danych pozwalających przeanalizować różnice ze względu na płeć,
- stereotypowe przedstawianie kobiet i mężczyzn w treściach generowanych przez modele,
- ryzyko dyskryminacji w zatrudnieniu, ochronie zdrowia, dostępie do usług i nadzorze,
- możliwość wykorzystywania AI do przemocy i nękania ze względu na płeć ([UNESCO Women for Ethical AI 2024](#)).

Jeżeli model regularnie przedstawia osoby określonej płci w węższym zestawie ról, kompetencji lub zawodów, może utrzymywać istniejące przekonania społeczne.

Nierówności płci

Raport UNESCO poświęcony AI i płci wskazuje kilka powiązanych obszarów nierówności.

1. mniejszy udział kobiet w tworzeniu technologii i podejmowaniu decyzji dotyczących AI
2. niedostatek danych pozwalających dokładnie analizować różnice ze względu na płeć
3. stereotypowe przedstawianie kobiet i mężczyzn w treściach generowanych przez systemy
4. ryzyko dyskryminacji w zatrudnieniu, ochronie zdrowia, dostępie do usług i nadzorze
5. możliwość wykorzystywania AI do przemocy i nękania ze względu na płeć

UNESCO Women for Ethical AI 2024

Jeżeli system regularnie przedstawia osoby określonej płci w węższym zestawie ról, kompetencji lub zawodów, może utrzymywać istniejące przekonania społeczne.

EDU x AI

5.2. Zagrożenia informacyjne

GenAI ułatwia tworzenie tekstów, obrazów, nagrań i syntetycznych głosów, które mogą wyglądać wiarygodnie niezależnie od ich zgodności z rzeczywistością. Zwiększa również skalę i tempo przygotowywania treści. Ocena zagrożenia wymaga jednak rozróżnienia kilku zjawisk. Nie każda fałszywa informacja jest dezinformacją, a treść prawdziwa także może zostać wykorzystana do wyrządzenia szkody.

Trzy rodzaje zagrożeń informacyjnych

Rodzaj zagrożenia zależy od dwóch cech: zgodności treści z rzeczywistością oraz intencji osoby, która ją tworzy lub rozpowszechnia.

W literaturze pojęcia te nie zawsze są stosowane konsekwentnie. W tym kursie przyjmujemy podział na misinformację, dezinformację i malinformację według [Wardle i Derakhshana \(2017\)](#), z uwzględnieniem ustaleń przeglądu systematycznego [Brody i Strömbäcka \(2024\)](#).

Malinformacja, dezinformacja, misinformacja

RODZAJ	CHARAKTER TREŚCI	INTENCJA	PRZYKŁAD
Malinformacja	treść prawdziwa, pozbawiona kontekstu albo ujawniona w szkodliwy sposób	wyrządzenie szkody	publikacja prywatnych danych lub fragmentu wypowiedzi zmieniającego jej sens
Dezinformacja	treść fałszywa, zmanipulowana albo wprowadzająca w błąd	wpływanie na emocje, opinie lub działania	sfabrykowana wypowiedź przypisana ekspertowi
Misinformacja	treść fałszywa lub błędnie opisana	brak intencji oszustwa	udostępnienie starego nagrania jako relacji z bieżącego wydarzenia

EDU x AI

1. Malinformacja: prawdziwa treść użyta w celu zaszkodzenia

Malinformacja to prawdziwa informacja wykorzystana w celu wyrządzenia szkody. Szkada może wynikać z ujawnienia treści prywatnej, usunięcia kontekstu albo przedstawienia autentycznego fragmentu w sposób zmieniający jego znaczenie. Ocena samej prawdziwości treści nie wystarcza, trzeba również sprawdzić jej pochodzenie, kontekst, kompletność i cel publikacji.

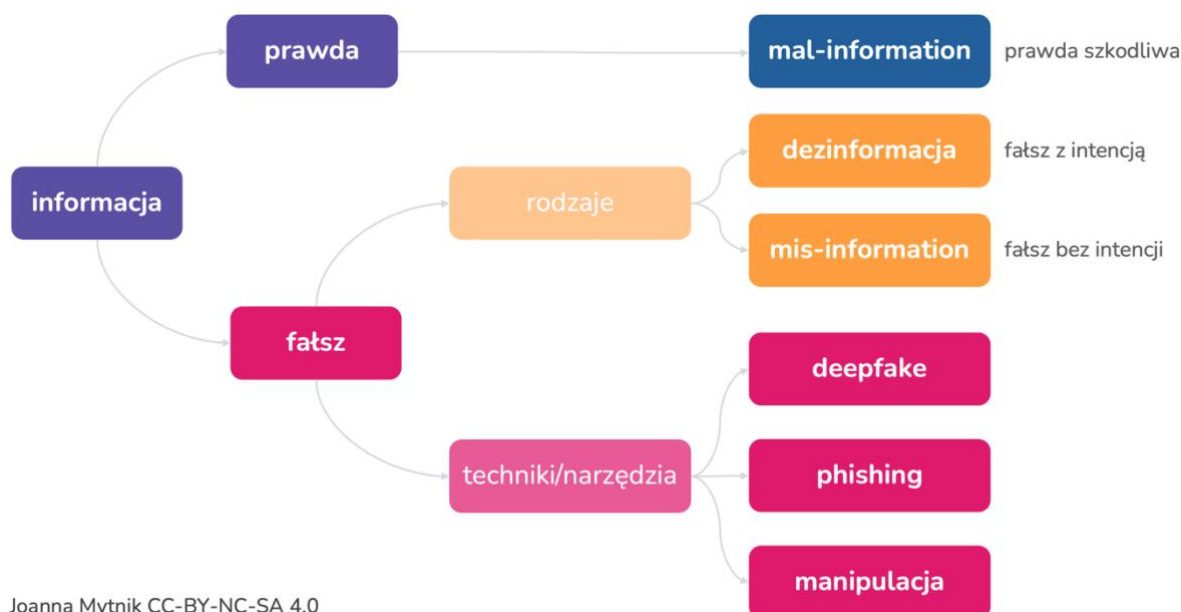
2. Dezinformacja: fałsz połączony z intencją manipulacji

Dezinformacja to fałszywa, zmanipulowana lub wprowadzająca w błąd treść świadomie tworzona albo rozpowszechniana w celu manipulacji lub wyrządzenia szkody. Jej celem może być wpływanie na opinie i zachowania, osłabienie zaufania, pogłębienie konfliktu lub destabilizacja.

3. Misinformacja: fałsz rozpowszechniany bez złej intencji

Misinformacja powstaje wtedy, gdy ktoś wierzy w prawdziwość błędnej treści i przekazuje ją dalej. Przykładami są stare nagranie przedstawione jako relacja z bieżącego wydarzenia, film udostępniony bez informacji o miejscu i czasie powstania, zdjęcie przypisane innemu zdarzeniu lub nieprawidłowo zinterpretowany wykres lub wynik badania.

Brak intencji nie usuwa konsekwencji. Wielokrotne udostępnianie błędnej informacji może zwiększać jej widoczność i wiarygodność społeczną.



Techniki wykorzystywane do oszustwa: deepfake

Deepfake jest syntetycznym materiałem wideo lub audio wygenerowanym z wykorzystaniem modeli GenAI. Technologia umożliwia między innymi podmianę twarzy, wygenerowanie głosu podobnego do głosu konkretnej osoby, umieszczenie osób w sytuacjach, które nie miały miejsca czy przygotowanie przekonującej wypowiedzi, której dana osoba nie wygłosiła.

Manipulacja może wykorzystywać prawdę i fałsz

Manipulacja polega na celowym wpływaniu na sposób postrzegania sytuacji lub podejmowania decyzji przez ukrywanie istotnych informacji, zniekształcanie treści albo wykorzystywanie emocji i podatności odbiorcy w sposób ograniczający możliwość świadomej oceny. Obejmuje techniki wpływania na odbiorcę bez jego świadomej zgody.

Wykorzystywanie języka moralno-emocjonalnego

Badania ([Brady i in. 2017](#)) pokazały związek między obecnością języka moralno-emocjonalnego a większą liczbą udostępnień w dyskusjach dotyczących polaryzujących tematów politycznych.

Clickbait to nagłówek lub treść zaprojektowana w celu przyciągnięcia uwagi i wywołania natychmiastowej reakcji. Wykorzystuje ciekawość, pobudzenie emocjonalne i poczucie pilności, w takich warunkach odbiorca może reagować przed sprawdzeniem źródła, kontekstu i dowodów.

Clickbait nie musi zawierać fałszu, ale może zniekształcać znaczenie materiału przez wyolbrzymienie, pominięcie warunków albo przedstawienie wniosku, którego nie uzasadnia dalsza treść.

O widoczności treści decydują nie tylko zachowania użytkowników, lecz także mechanizmy rekomendowania i porządkowania materiałów na platformach. Systemy optymalizowane pod kątem zaangażowania mogą zwiększać ekspozycję treści wywołujących silne reakcje.

Dlaczego fałsz może rozprzestrzeniać się szybciej?

[Vosoughi i in. \(2018\)](#) przeanalizowali około 126 000 historii rozpowszechnianych na Twitterze w latach 2006-2017. Obejmowały one ponad 4,5 miliona publikacji około trzech milionów użytkowników. Fałszywe wiadomości docierały dalej, szybciej, głębiej i szerzej niż wiadomości prawdziwe.

W analizowanym zbiorze:

- fałszywa wiadomość miała o 70% większe prawdopodobieństwo ponownego udostępnienia,
- prawdziwe informacje potrzebowały około sześciu razy więcej czasu, aby dotrzeć do 1500 osób,
- różnica była szczególnie wyraźna w przypadku informacji politycznych.

Autorzy wskazali również, że treści fałszywe były częściej nowe i zaskakujące. Ich przewaga w rozpowszechnianiu nie wynikała wyłącznie z działania kont automatycznych. Duże znaczenie miały decyzje ludzi o udostępnieniu materiału. Wynik tego badania dotyczy określonej platformy i okresu, dlatego nie należy automatycznie przenosić wartości liczbowych na wszystkie media i rodzaje informacji. Pokazuje jednak, że prawdziwa informacja nie uzyskuje przewagi wyłącznie dlatego, że jest zgodna z rzeczywistością.

Zanim zareagujesz lub udostępnisz

Gdy treść wywołuje silną emocję albo wymaga natychmiastowego działania, sprawdź:

- kto opublikował informację,
- czy źródło rzeczywiście istnieje,
- kiedy i w jakim kontekście powstał materiał,
- czy nagłówek odpowiada treści,
- czy dostępna jest pełna wypowiedź, nagranie lub dokument,
- czy informację potwierdzają niezależne źródła,
- czy komunikat nie zawiera prośby o dane, hasło, kod lub przelew,

Szczególnej ostrożności wymagają treści, które jednocześnie wywołują silne emocje, powołują się na autorytet i nakłaniają do szybkiego działania.

Zanim zareagujesz lub udostępnisz

Gdy treść wywołuje silną emocję albo wymaga natychmiastowego działania, sprawdź:

- 1 kto opublikował informację
- 2 czy źródło rzeczywiście istnieje
- 3 kiedy i w jakim kontekście powstał materiał
- 4 czy nagłówek odpowiada treści
- 5 czy dostępna jest pełna wypowiedź, nagranie lub dokument
- 6 czy informację potwierdzają niezależne źródła
- 7 czy komunikat nie zawiera prośby o dane, hasło, kod lub przelew

Szczegółowej ostrożności wymagają treści, które jednocześnie:

wywołują silne emocje

powołują się na autorytet

nakłaniają do szybkiego działania

EDU x AI

Kiedy błędna odpowiedź GenAI staje się zagrożeniem informacyjnym?

Błędna treść wygenerowana przez model nie jest sama w sobie dezinformacją, ponieważ technologia nie ma intencji wyrządzenia szkody ani manipulowania odbiorcą. Jeżeli użytkownik uzna ją za prawdziwą i rozpowszechni bez sprawdzenia, może stać się ona misinformacją.

Jeżeli człowiek celowo wykorzysta GenAI do przygotowania i rozpowszechniania fałszywych treści w celu manipulacji, mamy do czynienia z dezinformacją.

Do refleksji

Przypomnij sobie treść, którą ostatnio chciałeś lub chciałaś natychmiast udostępnić. Co wywołało tę reakcję: nowość, strach, oburzenie, ciekawość, autorytet nadawcy czy poczucie pilności?

Jakie informacje należało sprawdzić przed podjęciem działania?

5.3. Nierówny dostęp i nierówne korzyści

GenAI może przyczyniać się do zwiększenia możliwości edukacyjnych, technologia może **wspierać osoby neuroróżnorodne, obcojęzyczne, z różnymi potrzebami**, na przykład:

- upraszczać język materiałów,
- tłumaczyć teksty,
- tworzyć alternatywne sposoby wyjaśnienia,
- wspierać w komunikacji (zwłaszcza osoby w spektrum autyzmu),
- przekształcać tekst na mowę lub mowę na tekst,
- wspierać osoby uczące się w różnym tempie,
- ułatwiać rozpoczęcie pracy osobom doświadczającym barier językowych lub komunikacyjnych.

Korzyści nie rozkładają się jednak automatycznie i równo. Znaczenie mają:

- dostęp do płatnych wersji narzędzi,
- jakość połączenia internetowego i sprzętu,
- kompetencje cyfrowe oraz informacyjne,
- znajomość języka dominującego w danych treningowych,
- wiedza potrzebna do oceny odpowiedzi,
- dostępność narzędzia dla osób z różnymi potrzebami,
- możliwość odmowy korzystania z technologii bez utraty szans edukacyjnych.

[UNESCO](#) zwraca uwagę, że niewłaściwe wdrożenie GenAI może pogłębiać istniejące nierówności cyfrowe i edukacyjne. Zaleca zapewnienie równego dostępu, rozwijanie kompetencji oraz projektowanie rozwiązań skoncentrowanych na prawach i potrzebach człowieka ([Miao i Holmes 2023](#)).

Równy dostęp nie oznacza identycznego rozwiązania

Zapewnienie wszystkim tego samego narzędzia może być niewystarczające. Osoba posiadająca rozległą wiedzę w danej dziedzinie łatwiej rozpozna błąd i lepiej wykorzysta odpowiedź niż osoba początkująca. Dostęp do GenAI nie zastępuje więc dostępu do nauczyciela, dobrej jakości materiałów, informacji zwrotnej i warunków sprzyjających uczeniu się.

5.4. Koszt środowiskowy AI

Korzystanie z GenAI może sprawiać wrażenie działania niematerialnego: wpisujesz polecenie i po chwili otrzymujesz tekst, obraz albo kod. Za tą interakcją działa jednak infrastruktura obejmująca centra danych, procesory, sieci przesyłowe i systemy chłodzenia. Jej budowa i eksploatacja wymagają energii, wody oraz surowców.

Koszt środowiskowy GenAI nie powstaje wyłącznie w chwili generowania odpowiedzi. Obejmuje cały cykl życia technologii:

- wydobywanie i przetwarzanie surowców potrzebnych do produkcji sprzętu,
- produkcję procesorów, serwerów i infrastruktury sieciowej,
- trenowanie i dostrajanie modeli,
- codzienne wykonywanie zapytań, czyli inferencję,
- przechowywanie i przesyłanie danych,
- chłodzenie centrów danych,
- wymianę sprzętu oraz zagospodarowanie odpadów elektronicznych.

Nie istnieje jeden „koszt zapytania”

Koszt pojedynczej interakcji nie jest stałą właściwością GenAI. Zależy między innymi od:

- rodzaju, wielkości i architektury modelu,
- długości kontekstu oraz wygenerowanej odpowiedzi,
- rodzaju zadania i wymaganej liczby operacji,
- liczby generowanych wariantów,
- zastosowanego sprzętu i stopnia jego wykorzystania,
- dodatkowych operacji, takich jak wyszukiwanie, analiza plików lub generowanie obrazu,
- efektywności centrum danych i sposobu chłodzenia,
- źródeł energii dostępnych w danym miejscu i czasie.

Twierdzenia w rodzaju „jedno pytanie zużywa dokładnie określoną ilość energii lub wody” należy traktować ostrożnie. Wynik uzyskany dla jednego modelu, zadania i centrum danych nie może być automatycznie przeniesiony na wszystkie systemy.

Problemem pozostaje również **ograniczona przejrzystość dostawców**. Bez informacji o modelu, sprzęcie, miejscu wykonania obliczeń i przyjętej metodzie pomiaru trudno niezależnie zweryfikować podawane wartości.

AI w warunkach kryzysu klimatycznego

Koszt energetyczny AI należy rozpatrywać w warunkach trwającego kryzysu klimatycznego. Każda dodatkowa emisja gazów cieplarnianych przyczynia się do wzrostu ich stężenia w atmosferze i zwiększa skalę globalnego ocieplenia. Znaczenie ma zatem nie tylko ilość energii potrzebnej do działania systemów, lecz także źródło tej energii.

Ta sama operacja może mieć różny ślad węglowy w zależności od tego, czy została wykonana w centrum danych zasilanym energią ze źródeł niskoemisyjnych, czy energią

pochodzącą głównie z paliw kopalnych. Zakup energii ze źródeł odnawialnych również nie oznacza automatycznie, że w chwili wykonywania obliczeń do sieci nie trafiała energia z węgla lub gazu.

Dlatego ocenianie wpływu na podstawie samego zużycia energii albo deklaracji dostawcy może prowadzić do niepełnych wniosków.

Według [Międzynarodowej Agencji Energetycznej \(2025\)](#) zużycie energii elektrycznej przez centra danych może wzrosnąć do około 945 TWh w 2030 roku, czyli ponad dwukrotnie względem 2024 roku, AI ma być najważniejszym czynnikiem tego wzrostu. Jednak mimo dwukrotnego wzrostu centra danych to wciąż niecałe 3% globalnego zużycia energii elektrycznej w 2030 roku, a ich wkład w globalny wzrost popytu na energię jest mniejszy niż np. pojazdów elektrycznych czy klimatyzacji.

Prawie połowę dodatkowego zapotrzebowania mają pokryć źródła odnawialne, ale pozostała część będzie zaspokajana także przez gaz ziemny i węgiel. Rozwój infrastruktury AI może więc zwiększać emisje oraz konkurować o niskoemisyjną energię potrzebną do elektryfikacji transportu, ogrzewania i przemysłu.

Centra danych obsługują również wyszukiwarki, usługi chmurowe, transmisję multimedialną, handel elektroniczny i wiele innych. Nie zmienia to jednak faktu, że szybkie rozpowszechnienie AI zwiększa zapotrzebowanie na moc obliczeniową w okresie, w którym ograniczenie globalnych emisji wymaga równoczesnej dekarbonizacji wielu sektorów gospodarki.

Rodzaj zadania ma znaczenie

[Luccioni i in. \(2024\)](#) porównali zużycie energii przez modele wykonujące różne zadania, badanie wykazało znaczne różnice między zastosowaniami. Generowanie treści było zazwyczaj bardziej energochłonne niż ich klasyfikowanie, a generowanie obrazów wymagało znacznie więcej energii niż generowanie tekstu. W wielu zadaniach modele wielozadaniowe zużywały więcej energii niż mniejsze systemy przeznaczone do konkretnego zastosowania. Nie oznacza to, że każdy duży model będzie zawsze gorszym wyborem, model o większych możliwościach może wykonać zadanie, z którym mniejszy model sobie nie poradzi.

Badanie uzasadnia jednak ważne pytanie: **czy do danego zadania rzeczywiście potrzebny jest najbardziej rozbudowany dostępny system?**

Woda nie jest kosztem jednakowym w każdym miejscu

Woda może być używana bezpośrednio do chłodzenia centrów danych oraz pośrednio przy produkcji energii elektrycznej i sprzętu. Znaczenie tego zużycia zależy od miejsca i czasu. Litry wody wykorzystany w regionie zasobnym w wodę nie ma takich samych konsekwencji jak litr wykorzystany podczas suszy lub w miejscu, w którym o te same zasoby konkurują mieszkańcy, rolnictwo, przemysł i ekosystemy.

Dlatego sama objętość wody nie opisuje całego oddziaływania, potrzebna jest również informacja o jej źródle, lokalnych warunkach hydrologicznych, sposobie chłodzenia oraz o tym, czy podawana wartość oznacza pobór wody, czy jej rzeczywiste zużycie ([Li i in. 2025](#)).

Koszt środowiskowy infrastruktury cyfrowej ma zatem także wymiar przestrzenny i społeczny. Korzyści z użycia modelu mogą trafiać do użytkowników znajdujących się daleko od miejsca, w którym ponoszone są koszty energetyczne, wodne i infrastrukturalne.

Ślad materialny i odpady elektroniczne

AI to nie tylko energia, woda i emisja. Produkcja wyspecjalizowanych procesorów wymaga surowców, złożonych procesów przemysłowych oraz znacznych ilości energii i wody. Szybki rozwój modeli może również przyspieszać wymianę serwerów, mimo że starszy sprzęt pozostaje technicznie sprawny.

[Wang i in. \(2024\)](#) oszacowali, że skumulowana ilość odpadów elektronicznych związanych z rozwojem AI może w latach 2020-2030 wynieść od 1,2 do 5 milionów ton. Jest to prognoza zależna od przyjętych scenariuszy, a nie pomiar przyszłego wyniku. Autorzy pokazali jednak również, że wydłużanie życia sprzętu, jego ponowne wykorzystanie i inne strategie gospodarki obiegu zamkniętego mogą istotnie ograniczyć ten problem.

Nie tylko koszt, lecz także rezultat

Sama informacja o zużyciu energii nie rozstrzyga, czy dane zastosowanie jest uzasadnione. AI może wspierać badania naukowe, analizowanie danych, dostępność materiałów lub rozwiązywanie problemów środowiskowych. Może też służyć do masowego tworzenia treści o niewielkiej wartości, niepotrzebnych wariantów i materiałów, które nie zostaną wykorzystane.

Ocena zastosowania wymaga więc uwzględnienia co najmniej trzech elementów:

- zasobów potrzebnych do wykonania zadania,
- wartości uzyskanego rezultatu,
- dostępności mniej czasochłonnego sposobu osiągnięcia tego samego celu.

Warto unikać zastosowań, w których złożona infrastruktura zastępuje prostsze i wystarczające rozwiązanie.

Energia jądrowa dla infrastruktury AI

Rosnące zapotrzebowanie centrów danych na energię skłania firmy technologiczne do wspierania energetyki jądrowej, między innymi przez utrzymywanie istniejących elektrowni i inwestowanie w małe reaktory modułowe.

Energia jądrowa może ograniczyć emisje związane z zasilaniem infrastruktury AI i zapewnić jej stałe dostawy energii, ale nie zmniejsza liczby wykonywanych obliczeń ani całkowitego zużycia zasobów.

Według [IEA \(2025\)](#) energia jądrowa i odnawialna mogą w 2030 roku pokrywać niemal 60% zapotrzebowania centrów danych na energię. Projekty wspierane przez Microsoft, Google i Amazon pokazują jednak także skalę infrastruktury energetycznej uruchamianej w odpowiedzi na rozwój AI.

Koszt środowiskowy AI

Oddziaływanie GenAI obejmuje energię, wodę, emisje, surowce, produkcję sprzętu i odpady elektroniczne.

energia

woda

emisje

surowce

produkcja
sprzętu

odpady
elektroniczne

Nie istnieje jedna wiarygodna wartość opisująca koszt każdego zapytania

Rodzaj zadania, model, infrastruktura, miejsce i skala użycia mogą zasadniczo zmieniać wynik.

rodzaj zadania

model

infrastruktura

miejsce

skala użycia

Efektywność pojedynczej operacji to nie wszystko

Wzrost efektywności pojedynczej operacji nie musi prowadzić do zmniejszenia całkowitego zużycia zasobów.



niższy koszt
operacji



więcej
operacji

Odpowiedzialna decyzja polega na porównaniu kosztu, wartości rezultatu i dostępnych alternatyw. Zasada proporcjonalności nie oznacza rezygnacji z GenAI, lecz dobór rozwiązania adekwatnego do celu.

Wpływ GenAI na klimat zależy od skali obliczeń i źródeł energii, dlatego wzrost wykorzystania technologii należy oceniać w kontekście konieczności ograniczania globalnych emisji.

EDU x AI

Zasada proporcjonalności

Korzystaj z GenAI wtedy, gdy oczekiwana wartość rezultatu uzasadnia wykorzystanie zasobów. W praktyce:

- zacznij od określenia celu, a dopiero potem wybierz narzędzie,
- użyj prostszego modelu, jeżeli wystarczy do wykonania zadania,
- ogranicz długość kontekstu i odpowiedzi do tego, co rzeczywiście potrzebne,
- nie generuj wielu wariantów bez kryterium ich porównania,
- nie wykorzystuj generatora obrazów do tworzenia dekoracji bez znaczenia dla celu,
- poprawiaj polecenie na podstawie rozpoznanego problemu, zamiast wielokrotnie generować odpowiedź bez zmiany strategii,
- ponownie wykorzystaj sprawdzone materiały, jeżeli pozostają aktualne,
- oceniaj koszty także w skali całej grupy, przedmiotu i instytucji.

Zasada proporcjonalności

Korzystaj z GenAI wtedy, gdy oczekiwana wartość rezultatu uzasadnia wykorzystanie zasobów.

W praktyce:

1 zacznij od określenia celu, a dopiero potem wybierz narzędzie

2 użyj prostszego systemu, jeżeli wystarczy do wykonania zadania

3 ogranicz długość kontekstu i odpowiedzi do tego, co rzeczywiście potrzebne

4 nie generuj wielu wariantów bez kryterium ich porównania

5 nie wykorzystuj generatora obrazów do tworzenia dekoracji bez znaczenia dla celu

6 poprawiaj polecenie na podstawie rozpoznanego problemu, zamiast wielokrotnie generować odpowiedź bez zmiany strategii

7 ponownie wykorzystuj sprawdzone materiały, jeżeli pozostają aktualne

8 oceniaj koszty także w skali całej grupy, przedmiotu i instytucji

9 wymagając użycia konkretnego narzędzia, sprawdź, czy jego zastosowanie wnosi wartość niemożliwą do uzyskania prostszą metodą

EDU × AI

5.5. GenAI jako rozmówca i źródło wsparcia emocjonalnego

Model konwersacyjny może generować wypowiedzi **przypominające empatię**, zainteresowanie i troskę. Może zwracać się bezpośrednio do użytkownika, **dostosowywać styl** kolejnych odpowiedzi i podtrzymywać rozmowę.

Nie oznacza to, że model doświadcza emocji czy rozumie sytuację użytkownika. Odpowiedzi są generowane na podstawie przetwarzanych danych i wzorców językowych.

Wypowiedź przypominająca empatię to nadal generowanie tekstu

CO MODEL MOŻE ZROBIĆ

Model konwersacyjny może generować wypowiedzi przypominające empatię, zainteresowanie i troskę.

zwracać się bezpośrednio do użytkownika

dostosowywać styl kolejnych odpowiedzi

podtrzymywać rozmowę

CZEGO TO NIE OZNACZA

Nie oznacza to, że system doświadcza emocji, rozumie sytuację użytkownika albo tworzy relację w ludzkim znaczeniu.

doświadczenie emocji

rozumienie sytuacji użytkownika

relacja w ludzkim znaczeniu

! Odpowiedzi są generowane na podstawie przetwarzanych danych i wzorców językowych.

EDU x AI

Dlaczego interakcja może wydawać się relacją?

System jest zwykle:

- dostępny o każdej porze,
- szybki w udzielaniu odpowiedzi,
- skoncentrowany na treści wprowadzanej przez użytkownika,
- pozbawiony oznak zniecierpliwienia,
- zdolny do generowania potwierdzających i wspierających komunikatów,
- gotowy do kontynuowania rozmowy przez długi czas.

Te cechy mogą sprzyjać antropomorfizacji i przypisywaniu technologii intencji, troski lub wyjątkowej więzi.

Co mówią badania?

W eksperymentach dotyczących narzędzi zaprojektowanych jako wirtualni towarzysze krótkie interakcje mogły czasowo **zmniejszać deklarowane poczucie samotności**, szczególnie gdy użytkownik miał poczucie, że jego wypowiedź została zauważona i uwzględniona ([De Freitas i in. 2024](#)).

Jednocześnie badanie przekrojowe [Lai i in. \(2025\)](#) przeprowadzone wśród studentów wykazało **związek między objawami depresji, samotnością i wykorzystywaniem konwersacyjnej AI do towarzystwa**. Konstrukcja badania nie pozwala stwierdzić, czy korzystanie z technologii nasila trudności, czy osoby doświadczające trudności częściej szukają takiej formy kontaktu.

Korzyść odczuwana w danym momencie nie rozstrzyga więc o długoterminowym wpływie interakcji na relacje, zdrowie psychiczne i poszukiwanie profesjonalnej pomocy.

5.6. System konwersacyjny nie zastępuje profesjonalnej pomocy

Ogólnodostępny model językowy nie powinien być traktowany jako psycholog, psychoterapeuta, psychiatra, lekarz ani narzędzie diagnostyczne.

System może:

- nie rozpoznać sytuacji kryzysowej,
- zbyt łatwo potwierdzić założenie użytkownika,
- wygenerować nieadekwatną lub niebezpieczną poradę,
- pominąć ważne informacje,
- nie uwzględnić historii zdrowia i kontekstu społecznego.

W badaniu [Scholich i in. \(2025\)](#) porównującym odpowiedzi ogólnych modeli konwersacyjnych z odpowiedziami licencjonowanych terapeutów modele generowały odpowiedzi, które w części sytuacji **nie spełniały standardów** wymaganych w profesjonalnym wsparciu zdrowia psychicznego. Autorzy podkreślają potrzebę ostrożności, nadzoru i dalszych badań.

Specjalistyczne interwencje cyfrowe projektowane i badane do określonego zastosowania klinicznego są inną kategorią niż ogólny model konwersacyjny.

Ogólny model konwersacyjny a profesjonalne wsparcie

W badaniu porównującym odpowiedzi ogólnych modeli konwersacyjnych z odpowiedziami licencjonowanych terapeutów systemy generowały odpowiedzi, które w części sytuacji nie spełniały standardów wymaganych w profesjonalnym wsparciu zdrowia psychicznego.

Scholich i in. 2025

CZEGO POTRZEBA

Ostrożność, nadzór i dalsze badania

Autorzy podkreślają potrzebę ostrożności, nadzoru i dalszych badań.

WAŻNE ROZRÓŻNIENIE

To dwie różne kategorie

Specjalistyczne interwencje cyfrowe projektowane i badane do określonego zastosowania klinicznego są inną kategorią niż ogólny model konwersacyjny.

Sygnały ostrzegawcze

Warto zatrzymać interakcję i zwrócić się do człowieka, gdy:

- model językowy staje się głównym źródłem kontaktu i wsparcia emocjonalnego,
- korzystanie z niego zastępuje rozmowy z bliskimi lub specjalistą,
- użytkownik odczuwa silny niepokój, gdy narzędzie jest niedostępne,
- odpowiedzi modelu wpływają na istotne decyzje dotyczące zdrowia lub relacji,
- użytkownik zaczyna przypisywać modelowi świadomość, uczucia lub szczególną więź,
- model potwierdza przekonania, które budzą niepokój lub prowadzą do ryzykownych działań,
- rozmowa dotyczy samookaleczenia, przemocy, myśli samobójczych lub bezpośredniego zagrożenia.

Sygnały ostrzegawcze

Warto zatrzymać interakcję i zwrócić się do człowieka, gdy:

- ! model językowy staje się głównym źródłem kontaktu i wsparcia emocjonalnego
- ! korzystanie z niego zastępuje rozmowy z bliskimi lub specjalistą
- ! użytkownik odczuwa silny niepokój, gdy narzędzie jest niedostępne
- ! odpowiedzi modelu wpływają na istotne decyzje dotyczące zdrowia lub relacji
- ! użytkownik zaczyna przypisywać systemowi świadomość, uczucia lub szczególną więź
- ! model potwierdza przekonania, które budzą niepokój lub prowadzą do ryzykownych działań
- ! rozmowa dotyczy samookaleczenia, przemocy, myśli samobójczych lub bezpośredniego zagrożenia

Szukaj pomocy u specjalistów

W sytuacji kryzysu potrzebny jest kontakt z człowiekiem, odpowiednią instytucją i profesjonalną pomocą.

- **800 70 22 22** – Centrum Wsparcia dla Osób Dorosłych w Kryzysie Psychicznym (czynne 24/7, bezpłatne)
- **116 123** – Kryzysowy Telefon Zaufania dla dorosłych w kryzysie emocjonalnym (czynny 24/7, bezpłatny)
- **116 111** – Telefon zaufania dla dzieci i młodzieży (czynny 24/7, bezpłatny)
- **22 484 88 01** – Antydepresyjny Telefon Zaufania (Fundacja ITAKA)
- **112** – Numer alarmowy w sytuacji bezpośredniego zagrożenia życia lub zdrowia

Szukaj pomocy u specjalistów

800 70 22 22

Centrum Wsparcia dla Osób Dorosłych w Kryzysie Psychicznym
czynne 24/7, bezpłatne

116 123

Kryzysowy Telefon Zaufania dla dorosłych w kryzysie emocjonalnym
czynny 24/7, bezpłatny

116 111

Telefon zaufania dla dzieci i młodzieży
czynny 24/7, bezpłatny

22 484 88 01

Antydepresyjny Telefon Zaufania
Fundacja ITAKA

112

Numer alarmowy w sytuacji bezpośredniego zagrożenia życia lub zdrowia

5.7. Relacje są częścią uczenia się

Model może generować pytania, symulować stanowisko w dyskusji lub tworzyć propozycję informacji zwrotnej, ale nie uczestniczy w życiu społeczności akademickiej i nie ponosi wzajemnej odpowiedzialności za relacje oraz jej konsekwencje.

Jeżeli interakcja z GenAI zastępuje każdą konsultację, rozmowę zespołową i kontakt między studentem a nauczycielem, może ograniczać społeczny wymiar uczenia się.

Perspektywa studenta

Sprawdź:

- czy korzystanie z GenAI wspiera przygotowanie do rozmowy z ludźmi, czy tę rozmowę zastępuje,
- czy omawiasz swoje rozumowanie z innymi osobami,
- czy pytasz prowadzącego, gdy problem dotyczy wymagań przedmiotu,
- czy szukasz pomocy człowieka w sprawach dotyczących zdrowia, kryzysu lub ważnych decyzji życiowych.

Perspektywa nauczyciela akademickiego

Sprawdź:

- czy wprowadzenie GenAI pozostawia miejsce na dialog ze studentami,
- czy technologia nie zastępuje informacji zwrotnej wymagającej znajomości konkretnej osoby i sytuacji,
- czy student może uzyskać pomoc od człowieka,
- czy aktywność nadal wspiera współpracę i przynależność do społeczności,
- czy automatyzacja nie przenosi całej odpowiedzialności za uczenie się na studenta.

Świadome korzystanie z GenAI: pięć pytań

Pytanie 1. Komu może służyć rezultat?

Sprawdź, czy zastosowanie technologii przynosi rzeczywistą korzyść osobie uczącej się, nauczycielowi lub społeczności.

Pytanie 2. Kogo może pomijać lub przedstawiać stereotypowo?

Zwróć uwagę na reprezentację różnych grup, języków, perspektyw i doświadczeń.

Pytanie 3. Czy można sprawdzić informacje i ich źródła?

Nie wykorzystuj płynności językowej jako kryterium wiarygodności.

Pytanie 4. Czy korzyść uzasadnia wykorzystanie zasobów?

Rozważ prostsze narzędzie lub wykonanie zadania bez GenAI.

Pytanie 5. Czy technologia wspiera kontakt między ludźmi, czy zaczyna go zastępować?

Szczególną ostrożność zachowaj w sytuacjach dotyczących samotności, zdrowia psychicznego, konfliktu i kryzysu.

KARTA PRACY

Świadome korzystanie z GenAI: pięć pytań

1 Komu może służyć rezultat?

Sprawdź, czy zastosowanie technologii przynosi rzeczywistą korzyść osobie uczącej się, nauczycielowi lub społeczności.

2 Kogo może pomijać lub przedstawiać stereotypowo?

Zwróć uwagę na reprezentację różnych grup, języków, perspektyw i doświadczeń.

3 Czy można sprawdzić informacje i ich źródła?

Nie wykorzystuj płynności językowej jako kryterium wiarygodności.

4 Czy korzyść uzasadnia wykorzystanie zasobów?

Rozważ prostsze narzędzie lub wykonanie zadania bez GenAI.

5 Czy technologia wspiera kontakt między ludźmi, czy zaczyna go zastępować?

Szczególne ostrożność zachowaj w sytuacjach dotyczących samotności, zdrowia psychicznego, konfliktu i kryzysu.

Podsumowanie modułu 5

Moduł 5 OSIEM WNIOSKÓW

Podsumowanie

— Skąd biorą się treści

- 1 Systemy GenAI odzwierciedlają dane, decyzje projektowe i warunki społeczne, w których zostały stworzone oraz wdrożone.
- 2 Treści generowane przez modele mogą powielać stereotypy i nierówno reprezentować różne grupy oraz perspektywy.
- 3 Naturalny styl i spójność wypowiedzi nie są dowodami prawdziwości informacji.

— Koszty i dostęp

- 4 GenAI może wspierać dostępność edukacji, ale nierówny dostęp do narzędzi, wiedzy i kompetencji może pogłębiać istniejące różnice.
- 5 Trenowanie i używanie modeli wymaga zasobów, dlatego warto oceniać proporcję między korzyścią a kosztem.

energia

woda

infrastruktura

sprzęt

— Relacje i wsparcie

- 6 Interakcja z systemem może chwilowo zmniejszać odczuwane poczucie samotności, ale długoterminowe skutki zależą od sposobu i intensywności korzystania oraz sytuacji użytkownika.
- 7 Ogólny model konwersacyjny nie zastępuje relacji międzyludzkiej ani profesjonalnego wsparcia psychologicznego i medycznego.

- 8 Odpowiedzialna decyzja o użyciu GenAI uwzględnia nie tylko jakość rezultatu, ale także wpływ na:

ludzi

relacje

dostęp do
edukacji

środowisko

obieg informacji

Moduł 6. Studiowanie z GenAI: jak korzystać z technologii i zachować samodzielność

Czas realizacji: ok. 50 minut

W tym module nauczysz się dobierać sposób wykorzystania GenAI do celu zadania, chronić własną samodzielność oraz sprawdzać, co potrafisz wyjaśnić i wykonać bez dostępu do technologii.

Pytanie otwierające

Wyobraź sobie, że dzięki GenAI wykonujesz zadania szybciej i otrzymujesz lepsze oceny. Jak sprawdzisz, czy technologia rzeczywiście wspiera rozwój Twoich kompetencji, a nie tylko poprawia jakość oddawanych prac?

— PYTANIE OTWIERAJĄCE



Wyobraź sobie, że dzięki GenAI wykonujesz zadania szybciej i otrzymujesz lepsze oceny. Jak sprawdzisz, czy technologia rzeczywiście wspiera rozwój Twoich kompetencji, a nie tylko poprawia jakość oddawanych prac?

6.1. Sprawne korzystanie z narzędzia nie jest jeszcze kompetencją AI

Umiejętność wpisania polecenia i uzyskania użytecznej odpowiedzi jest tylko jednym z elementów kompetencji AI.

[Long i Magerko \(2020\)](#) definiują kompetencje AI jako zestaw umiejętności pozwalających krytycznie oceniać technologie AI, skutecznie się z nimi komunikować i wykorzystywać je w różnych obszarach życia. [Ng i współautorzy \(2021\)](#) wyróżniają cztery powiązane obszary kompetencji:

1. rozumienie podstaw działania AI,
2. wykorzystywanie technologii,
3. ocenianie jej rezultatów,
4. rozpoznawanie konsekwencji etycznych.

UNESCO dodaje do tego perspektywę skoncentrowaną na człowieku. Student powinien nie tylko korzystać z narzędzi, ale także oceniać ich wpływ na ludzi, społeczeństwo i środowisko oraz zachowywać kontrolę nad podejmowanymi decyzjami ([Miao i Shiohira 2024](#)).

Osoba posiadająca kompetencje cyfrowe (AI literacy):

- rozumie, że model generuje odpowiedzi na podstawie wzorców w danych,
- zna ograniczenia wykorzystywanego narzędzia,
- potrafi dobrać narzędzie do zadania,
- wie, kiedy GenAI nie jest potrzebna,
- potrafi sformułować użyteczne polecenie,
- sprawdza odpowiedzi w innych źródłach,
- zauważa błędy, uprzedzenia i pominięte perspektywy,
- chroni dane i respektuje zasady obowiązujące na uczelni,
- ujawnia zakres wykorzystania GenAI,
- potrafi pracować również bez dostępu do technologii.

Niewątpliwie dziś jedną z najważniejszych kompetencji jest **podejmowanie trafnych decyzji** dotyczących tego, kiedy, po co i w jaki sposób **wykorzystać GenAI**.

Kompetencje AI to więcej niż napisanie promptu

Umiejętność wpisania polecenia i uzyskania użytecznej odpowiedzi jest tylko jednym z elementów kompetencji AI. Long i Magerko (2020) definiują je jako zestaw umiejętności pozwalających krytycznie oceniać technologie AI, skutecznie się z nimi komunikować i wykorzystywać je w różnych obszarach życia.

Cztery powiązane obszary kompetencji Ng i współautorzy 2021

Rozumienie podstaw działania AI

Wykorzystywanie technologii

Ocenianie jej rezultatów

Rozpoznawanie konsekwencji etycznych

UNESCO dodaje perspektywę skoncentrowaną na człowieku. Student powinien nie tylko korzystać z narzędzi, ale także oceniać ich wpływ na ludzi, społeczeństwo i środowisko oraz zachowywać kontrolę nad podejmowanymi decyzjami.

Miao i Shiohira 2024

Osoba posiadająca kompetencje cyfrowe AI literacy

- ✓ rozumie, że model generuje odpowiedzi na podstawie wzorców w danych
- ✓ zna ograniczenia wykorzystywanego narzędzia
- ✓ potrafi dobrać narzędzie do zadania
- ✓ wie, kiedy GenAI nie jest potrzebna
- ✓ potrafi sformułować użyteczne polecenie
- ✓ sprawdza odpowiedzi w innych źródłach
- ✓ zauważa błędy, uprzedzenia i pominięte perspektywy
- ✓ chroni dane i respektuje zasady obowiązujące na uczelni
- ✓ ujawnia zakres wykorzystania GenAI
- ✓ potrafi pracować również bez dostępu do technologii

Najważniejszą kompetencją jest podejmowanie trafnych decyzji dotyczących tego, kiedy, po co i w jaki sposób wykorzystać GenAI.

6.2. Mapa odpowiedzialności

Przed rozpoczęciem zadania podziel czynności na trzy grupy.

Muszę wykonać samodzielnie

Są to czynności związane bezpośrednio z **celem** uczenia się, na przykład:

- zdefiniowanie problemu,
- wybór metody,
- analiza danych,
- ocena jakości źródeł,
- interpretacja wyników,
- uzasadnienie decyzji,
- sformułowanie wniosków.

Mogę wykonać przy wsparciu GenAI

Są to czynności, podczas których technologia dostarcza **rusztowania**, ale nie przejmuje odpowiedzialności, na przykład:

- zadawanie pytań naprowadzających,
- wskazywanie luk w wyjaśnieniu,
- generowanie dodatkowych przykładów,
- udzielanie informacji zwrotnej,
- sprawdzanie rozumienia,
- proponowanie przypadków testowych.

Mogę delegować, jeśli pozwalają na to zasady

Są to czynności, których samodzielne wykonanie nie jest celem zadania, na przykład:

- transkrypcja,
- wstępne formatowanie,
- korekta językowa,
- przekształcenie danych do wskazanego formatu,
- przygotowanie technicznej wersji materiału.

Mapa odpowiedzialności zależy od zadania. W jednym przedmiocie napisanie kodu może być podstawowym efektem uczenia się, a w innym kod może być jedynie narzędziem potrzebnym do analizy danych.

Mapa odpowiedzialności

Przed rozpoczęciem zadania podziel czynności na trzy grupy.



EDU x AI

6.3. Czytanie, praca ze źródłami, streszczanie

GenAI może wyjaśnić trudny fragment, zaproponować strukturę notatki lub wygenerować pytania do tekstu. **Streszczenie nie jest jednak neutralną, skróconą kopią artykułu.**

Jego przygotowanie wymaga wyboru informacji uznanych za najważniejsze, pominięcia pozostałych oraz określenia relacji między nimi. Model może nadać różną wagę poszczególnym fragmentom, pominąć warunki ograniczające wnioski, zmienić poziom pewności twierdzenia albo przedstawić wynik jako bardziej ogólny, niż uzasadnia to badanie.

[Peters i Chin-Yee \(2025\)](#) porównali 4900 streszczeń wygenerowanych przez dziesięć modeli ze streszczanymi tekstami naukowymi. W zależności od modelu od 26% do 73% streszczeń **nadmiernie uogólniało wyniki**, a streszczenia modeli były niemal pięć razy częściej nadmiernie uogólnione niż streszczenia przygotowane przez ludzi.

[Tang i in. \(2023\)](#) wykazali natomiast, że podczas streszczania dowodów medycznych modele **pomijały istotne cechy badań, zmieniały poziom pewności wniosków i popełniały błędy** interpretacyjne. Trudności nasilały się przy dłuższych tekstach.

Ryzykowne jest zatem traktowanie wygenerowanego streszczenia jako substytutu lektury.

W badaniu obejmującym 195 osób w wieku 18-22 lat przeczytanie streszczenia wygenerowanego przez model zamiast pełnego tekstu pogorszyło rozumienie w grupie osób o wyższych wyjściowych kompetencjach czytania. W przypadku osób z niższymi wynikami efekty nie były jednoznaczne ([Etkin i in. 2025](#)).

Badanie dotyczyło tekstów standaryzowanego testu, a nie artykułów naukowych, dlatego jego wyników nie należy przenosić bezpośrednio na każdą sytuację akademicką. Pokazuje ono jednak, że zwięzłość uzyskana dzięki streszczeniu może oznaczać utratę szczegółów i niuansów potrzebnych do pełnego rozumienia.

Streszczenie wygenerowane przed lekturą może dodatkowo stworzyć wstępną ramę interpretacyjną: wskazać, które problemy wydają się centralne, a które drugorzędne. Jest to prawdopodobny mechanizm, ale wymaga jeszcze bezpośredniego zbadania. Dlatego bezpieczniej wykorzystać GenAI po pierwszym kontakcie z tekstem, na przykład do porównania własnego streszczenia, wykrywania luk w notatkach, formułowania pytań albo wyjaśnienia konkretnego fragmentu.

Czytanie, praca ze źródłami, streszczanie

GenAI może wyjaśnić trudny fragment, zaproponować strukturę notatki lub wygenerować pytania do tekstu. Streszczenie nie jest jednak neutralną, skróconą kopią artykułu.

Przygotowanie streszczenia wymaga decyzji

wyboru informacji uznanych za najważniejsze

pominięcia pozostałych

określenia relacji między nimi

Co może się zmienić w wygenerowanym streszczeniu

! model może nadać różną wagę poszczególnym fragmentom

! może pominąć warunki ograniczające wnioski

! może zmienić poziom pewności twierdzenia

! może przedstawić wynik jako bardziej ogólny, niż uzasadnia to badanie

Zwięzłość uzyskana dzięki streszczeniu może oznaczać utratę szczegółów i niuansów potrzebnych do pełnego rozumienia.

Traktowanie wygenerowanego streszczenia jako substytutu lektury jest ryzykowne.

Prompt wspierający analizę tekstu

Pracuj wyłącznie na treści załączonego artykułu. Pomóż mi zrozumieć wskazany fragment. Najpierw wyjaśnij użyte pojęcia, następnie pokaż związek między tym fragmentem a pytaniem badawczym. Oddziel informacje zawarte w artykule od własnych objaśnień. Jeżeli czegoś nie ma w tekście, zaznacz brak informacji.

Czego nie delegować?

Nie deleguj całej lektury, jeżeli masz:

- ocenić metodę badania,
- porównać wyniki kilku publikacji,
- rozpoznać ograniczenia,
- sformułować stanowisko,
- zdecydować, czy źródło wspiera Twoje twierdzenie.

6.4. Pisanie pracy

GenAI może ułatwić rozpoczęcie pisania, ale wygenerowany tekst może narzucić strukturę argumentacji, zanim samodzielnie określisz własne stanowisko.

Bezpieczniejsza kolejność

Najpierw samodzielnie:

1. określ pytanie lub problem,
2. zapisz główną tezę,
3. wybierz źródła,
4. przygotuj mapę argumentów,
5. określ, jaki dowód wspiera każdy argument,
6. napisz pierwszą wersję kluczowego fragmentu.

Następnie wykorzystaj GenAI do:

- wygenerowania kontrargumentów,
- wskazania niejasnych przejść,
- sprawdzenia spójności struktury,
- przygotowania pytań do tekstu,
- identyfikacji twierdzeń wymagających źródła,
- korekty językowej w zakresie dozwolonym przez prowadzącego.

Pisanie pracy

GenAI może ułatwić rozpoczęcie pisania, ale wygenerowany tekst może narzucić strukturę argumentacji, zanim samodzielnie określisz własne stanowisko.

Bezpieczniejsza kolejność



EDU x AI

Prompt dający informację zwrotną

Przeanalizuj mój tekst na podstawie podanych kryteriów. Nie przepisuj go i nie dodawaj nowych argumentów. Wskaż miejsca, w których związek między twierdzeniem a dowodem jest niejasny. W każdym takim przypadku zadaj pytanie, które pomoże mi samodzielnie poprawić argumentację.

Sprawdź autorstwo własnej pracy

Przed oddaniem tekstu odpowiedz:

- czy główna teza jest moja,
- czy samodzielnie wybrałam/em źródła,
- czy przeczytałam/em cytowane publikacje,
- czy rozumiem każde użyte pojęcie,
- czy potrafię wyjaśnić strukturę argumentacji,
- czy zgadzam się z każdym wnioskiem,
- czy ujawniłam/em rzeczywisty zakres wykorzystania GenAI.

6.5. Rozwiązywanie zadań

Największe ryzyko pojawia się wtedy, gdy model generuje poprawnie wyglądające rozwiązanie zanim podejmiesz własną próbę. Gotowa odpowiedź może wydawać się zrozumiała, ponieważ łatwo śledzić kolejne kroki. Nie oznacza to, że potrafisz samodzielnie wybrać i zastosować tę metodę.

Sekwencja pracy

1. Zapisz dane i niewiadome.
2. Określ, czego dotyczy problem.
3. Przywołaj możliwe zasady, wzory lub metody.
4. Wybierz metodę i uzasadnij wybór.
5. Podejmij własną próbę.
6. Dopiero wtedy skorzystaj z podpowiedzi.
7. Popraw rozwiązanie.
8. Rozwiąż podobne zadanie bez GenAI.

Rozwiązywanie zadań

Największe ryzyko pojawia się wtedy, gdy model generuje poprawnie wyglądające rozwiązanie zanim podejmiesz własną próbę. Gotowa odpowiedź może wydawać się zrozumiała, ponieważ łatwo śledzić kolejne kroki. Nie oznacza to, że potrafisz samodzielnie wybrać i zastosować tę metodę.

Sekwencja pracy

1 Zapisz dane i niewiadome.

2 Określ, czego dotyczy problem.

3 Przywołaj możliwe zasady, wzory lub metody.

4 Wybierz metodę i uzasadnij wybór.

5 Podejmij własną próbę.

6 Dopiero wtedy skorzystaj z podpowiedzi.

7 Popraw rozwiązanie.

8 Rozwiąż podobne zadanie bez GenAI.

EDU x AI

Prompt pracy etapowej

Nie rozwiązuj zadania. Przeanalizuj moją próbę i wskaż pierwszy etap, który wymaga ponownego sprawdzenia. Nie podawaj poprawnego działania. Zadaj jedno pytanie naprowadzające i poczekaj na moją odpowiedź.

Test transferu wiedzy

Po zakończeniu pracy zmień:

- dane liczbowe,
- warunki początkowe,
- reprezentację problemu,
- kontekst zastosowania,
- jedno z ograniczeń.

Następnie wykonaj nowe zadanie bez pomocy GenAI. Dopiero taki sprawdzian pokazuje, czy potrafisz zastosować metodę poza przykładem, nad którym pracowała technologia.

6.6. Programowanie

Wygenerowany kod może działać, mimo że użytkownik nie rozumie jego struktury. Może również zawierać błędy, luki bezpieczeństwa, niepotrzebne zależności lub rozwiązania niedopasowane do wymagań zadania.

Przed użyciem GenAI

Samodzielnie:

1. określ wymagania,
2. zapisz dane wejściowe i oczekiwany wynik,
3. rozplanuj algorytm,
4. wskaż przypadki brzegowe,
5. przygotuj własną próbę.

Podczas pracy

Możesz poprosić model o:

- wyjaśnienie komunikatu o błędzie,
- wskazanie kategorii problemu,
- wygenerowanie przypadków testowych,
- porównanie złożoności dwóch rozwiązań,
- wskazanie potencjalnych zagrożeń,
- zadawanie pytań prowadzących do znalezienia błędu.

Po wygenerowaniu kodu

Sprawdź:

- czy rozumiesz każdą funkcję i bibliotekę,
- czy kod spełnia wszystkie wymagania,
- czy działa dla przypadków typowych i brzegowych,
- czy poprawnie obsługuje błędy,
- czy nie ujawnia danych lub kluczy dostępu,
- czy użyte biblioteki i funkcje rzeczywiście istnieją,
- czy potrafisz odtworzyć najważniejsze elementy bez GenAI.

Jeżeli nie potrafisz wyjaśnić fragmentu kodu, nie traktuj go jako ukończonego elementu własnej pracy.

Programowanie

Wygenerowany kod może działać, mimo że użytkownik nie rozumie jego struktury. Może również zawierać błędy, luki bezpieczeństwa, niepotrzebne zależności lub rozwiązania niedopasowane do wymagań zadania.

1 Przed użyciem GenAI

samodzielnie

określ wymagania

zapisz dane wejściowe i oczekiwany wynik

rozplanuj algorytm

wskaż przypadki brzegowe

przygotuj własną próbę

2 Podczas pracy

możesz poprosić model o

wyjaśnienie komunikatu o błędzie

wskazanie kategorii problemu

wygenerowanie przypadków testowych

porównanie złożoności dwóch rozwiązań

wskazanie potencjalnych zagrożeń

zadawanie pytań prowadzących do znalezienia błędu

3 Po wygenerowaniu kodu, sprawdź

czy rozumiesz każdą funkcję i bibliotekę

czy kod spełnia wszystkie wymagania

czy działa dla przypadków typowych i brzegowych

czy poprawnie obsługuje błędy

czy nie ujawnia danych lub kluczy dostępu

czy użyte biblioteki i funkcje rzeczywiście istnieją

czy potrafisz odtworzyć najważniejsze elementy bez GenAI

Jeżeli nie potrafisz wyjaśnić fragmentu kodu, nie traktuj go jako ukończonego elementu własnej pracy.

6.7. Analiza danych

GenAI może zaproponować metodę analizy lub wygenerować kod, ale nie zna automatycznie znaczenia danych, warunków ich zebrania ani ograniczeń badania.

Przed wykorzystaniem narzędzia:

1. sprawdź, czy możesz udostępnić dane,
2. usuń dane osobowe i informacje poufne,
3. określ pytanie badawcze,
4. rozpoznaj rodzaj zmiennych,
5. zapisz założenia planowanej analizy,
6. przewidź, jakiego rodzaju wynik może odpowiedzieć na pytanie.

Po otrzymaniu propozycji:

- sprawdź założenia metody,
- oceń, czy metoda odpowiada strukturze danych,
- uruchom i przetestuj kod,
- porównaj wynik z innym sposobem analizy,
- samodzielnie zinterpretuj rezultat,
- oddziel wynik statystyczny od wniosku merytorycznego.

Wprowadzenie tabeli bez nazwisk nie zawsze oznacza, że dane są anonimowe. Zestaw rzadkich cech może umożliwić identyfikację osoby.

6.8. Praca zespołowa

W projekcie zespołowym GenAI może przyspieszyć przygotowanie materiałów, ale może również ukryć nierówny podział pracy i brak wspólnego rozumienia rezultatu.

W zespole wspólnie powinniście ustalić:

- do których czynności można wykorzystać GenAI,
- które elementy każdy członek zespołu wykonuje samodzielnie,
- jak będą sprawdzane wygenerowane treści,
- kto odpowiada za poszczególne decyzje,
- jak zostanie udokumentowany zakres wykorzystania technologii,
- czy każda osoba rozumie końcowy rezultat.

Test wspólnej odpowiedzialności. Każda osoba w zespole powinna potrafić:

- wyjaśnić główny problem,
- uzasadnić przyjętą metodę,
- opisać najważniejsze decyzje,
- wskazać ograniczenia rozwiązania,
- określić, gdzie wykorzystano GenAI,
- odpowiedzieć na pytania dotyczące własnego wkładu.

Wspólna prezentacja nie jest dowodem wspólnego uczenia się, jeżeli poszczególne osoby znają wyłącznie swój niewielki fragment pracy.

6.9. Samotestowanie z GenAI

Samo czytanie materiału może wywoływać poczucie znajomości, które łatwo pomylić z opanowaniem treści (iluzja kompetencji). Samotestowanie wymaga **wydobycia wiedzy z pamięci, zastosowania** jej i **sprawdzenia**, czego jeszcze nie potrafisz zrobić.

GenAI może generować pytania, prowadzić kolejne próby i przekazywać informację zwrotną, ale warunkiem jest **odpowiednia kolejność pracy** (jeżeli odpowiedź lub rozbudowana podpowiedź pojawi się przed Twoją próbą, zadanie przestaje sprawdzać samodzielne rozumienie).

1. Określ zakres i przygotuj materiały źródłowe

Do przygotowania pytań wykorzystaj przede wszystkim:

1. efekty uczenia się i zakres wymagań,
2. sylabus (kartę przedmiotu),
3. materiały udostępnione przez nauczyciela,
4. informacje zamieszczone na platformie eNauczanie PG,
5. polecane podręczniki i publikacje,
6. przykładowe zadania i problemy,
7. kryteria oceniania.

Nie wprowadzaj do zewnętrznego systemu danych osobowych, materiałów poufnych ani treści, których nie wolno udostępniać. Jeżeli nie możesz przekazać pełnego materiału, przygotuj własną notatkę zawierającą pojęcia, zależności i wymagania.

2. Ogranicz odpowiedzi do dostarczonych materiałów

„Korzystaj wyłącznie z załączonych materiałów, nie uzupełniaj odpowiedzi informacjami spoza nich. Jeżeli materiał nie wystarcza do przygotowania pytania lub oceny odpowiedzi, napisz to wprost. Każdą informację zwrotną powiąż z odpowiednim fragmentem materiału”.

3. Zaczynaj od wydobycia wiedzy z pamięci (retrieval practice)

Pierwsza próba powinna odbywać się **bez materiałów i bez podpowiedzi**. Najbardziej diagnostyczne są pytania otwarte, ponieważ wymagają samodzielnego sformułowania odpowiedzi. Rozpoznanie poprawnej opcji w quizie ABCD może być łatwiejsze niż samodzielne odtworzenie wiedzy.

Po każdej odpowiedzi oceń swoją pewność, na przykład w skali od 0 do 100%. Pozwala to rozróżnić:

- odpowiedź poprawną i pewną,
- odpowiedź poprawną, ale niepewną,
- odpowiedź błędną mimo dużej pewności,
- brak wiedzy rozpoznany przez Ciebie.

Szczególnie ważne są **błędne odpowiedzi udzielone z dużą pewnością**. Wskazują na nie brak pamięci, lecz nieprawidłowe rozumienie wymagające korekty.

RETRIEVAL PRACTICE

Zacznij od wydobycia wiedzy z pamięci

Pierwsza próba powinna odbywać się bez materiałów i bez odpowiedzi.

NAJBARDZIEJ DIAGNOSTYCZNE

Pytania otwarte

Wymagają samodzielnego sformułowania odpowiedzi.

MNIEJ WYMAGAJĄCE

Quiz ABCD

Rozpoznanie poprawnej opcji może być łatwiejsze niż samodzielne odtworzenie wiedzy.

Po każdej odpowiedzi oceń swoją pewność

na przykład w skali od 0 do 100%, co pozwala rozróżnić cztery przypadki

1 odpowiedź poprawna i pewna

2 odpowiedź poprawna, ale niepewna

3 odpowiedź błędna mimo dużej pewności

4 brak wiedzy rozpoznany przez Ciebie

Szczególnie ważne są błędne odpowiedzi udzielone z dużą pewnością. Wskazują one nie na brak pamięci, lecz na nieprawidłowe rozumienie wymagające korekty.

EDU x AI

4. Proś o informację zwrotną, nie tylko ocenę

Informacja „dobrze” albo „źle” nie wystarcza do dalszego uczenia się. Informacja zwrotna powinna wskazać:

- co w odpowiedzi jest poprawne,
- czego brakuje,
- który element jest błędny lub nieprecyzyjny,
- na czym polega błąd rozumowania,
- gdzie można sprawdzić odpowiednią informację,
- co należy zrobić w kolejnej próbie.

Po otrzymaniu informacji zwrotnej nie przechodź od razu do następnego pytania. Najpierw popraw odpowiedź własnymi słowami, następnie rozwiąż nowe zadanie sprawdzające ten sam mechanizm w innym kontekście. Przenoszenie wiedzy na nowe sytuacje nie zachodzi automatycznie, ale można je rozwijać przez odpowiednio zróżnicowane zadania.

Proś o informację zwrotną, nie tylko ocenę

Informacja „dobrze” albo „źle” nie wystarcza do dalszego uczenia się.

Informacja zwrotna powinna wskazać

- | | |
|--|--------------------------------------|
| 1 co w odpowiedzi jest poprawne | 2 czego brakuje |
| 3 który element jest błędny lub nieprecyzyjny | 4 na czym polega błąd rozumowania |
| 5 gdzie można sprawdzić odpowiednią informację | 6 co należy zrobić w kolejnej próbie |

Po otrzymaniu informacji zwrotnej nie przechodź od razu do następnego pytania.

NAJPIERW

popraw odpowiedź własnymi słowami



NASTĘPNIE

rozwiąż nowe zadanie sprawdzające ten sam mechanizm w innym kontekście

Przenoszenie wiedzy na nowe sytuacje nie zachodzi automatycznie, ale można je rozwijać przez odpowiednio zróżnicowane zadania.

EDU x AI

5. Stosuj różne rodzaje pytań

Dobra sesja nie powinna składać się wyłącznie z pytań o definicje. Poproś o zadania wymagające:

- odtworzenia pojęcia lub zależności z pamięci,
- wyjaśnienia mechanizmu własnymi słowami,
- zastosowania pojęcia w nowym kontekście,
- dobrania metody i uzasadnienia wyboru,
- znalezienia błędu w rozwiązaniu, kodzie, obliczeniu lub rozumowaniu,
- porównania dwóch rozwiązań,
- przewidzenia wyniku po zmianie jednego z warunków,
- interpretacji danych, wykresu lub schematu,
- sformułowania kontrprzykładu,
- ustnej obrony decyzji.

Pytania prawda/fałsz i ABCD mogą być przydatne, ale wymagają dobrze przygotowanych **dystraktorów** i uzasadnienia wyboru. Nie zatrzymuj się na wskazaniu odpowiedzi.

Wyjaśnij, dlaczego ją wybierasz i w jakich warunkach mogłaby przestać być poprawna.

Stosuj różne rodzaje pytań

Dobra sesja nie powinna składać się wyłącznie z pytań o definicje. Poproś o zadania wymagające:

1 odtworzenia pojęcia lub zależności z pamięci

2 wyjaśnienia mechanizmu własnymi słowami

3 zastosowania pojęcia w nowym kontekście

4 dobrania metody i uzasadnienia wyboru

5 znalezienia błędu w rozwiązaniu, kodzie, obliczeniu lub rozumowaniu

6 porównania dwóch rozwiązań

7 przewidzenia wyniku po zmianie jednego z warunków

8 interpretacji danych, wykresu lub schematu

9 sformułowania kontrprzykładu

10 ustnej obrony decyzji

Pytania prawda i fałsz oraz ABCD mogą być przydatne, ale wymagają dobrze przygotowanych dystraktorów i uzasadnienia wyboru.

Nie zatrzymuj się na wskazaniu odpowiedzi. Wyjaśnij, dlaczego ją wybierasz i w jakich warunkach mogłaby przestać być poprawna.

EDU x AI

6. Dostosuj trudność, ale nie usuwaj wysiłku

Jeżeli nie potrafisz odpowiedzieć, nie proś od razu o rozwiązanie. Zastosuj stopniowanie wsparcia:

- ponów samodzielną próbę,
- poproś o wskazanie luki lub rodzaju błędu,
- poproś o jedno pytanie naprowadzające,
- poproś o przypomnienie właściwego pojęcia,
- dopiero na końcu przeanalizuj pełne rozwiązanie.

Po skorzystaniu ze wskazówki wykonaj podobne zadanie bez wsparcia. Prawidłowa odpowiedź z podpowiedzią nie jest jeszcze dowodem samodzielnego opanowania umiejętności.

7. Prowadź rejestr błędów

Po sesji zapisz:

- czego dotyczyło pytanie,
- jakiej odpowiedzi udzieliłeś,
- na czym polegał błąd,
- dlaczego go popełniłeś,
- jaka zasada pozwala go skorygować,
- kiedy ponownie sprawdzisz ten obszar.

Nie zapisuj wyłącznie poprawnej odpowiedzi. Nazwij mechanizm błędu, na przykład: pomylenie dwóch pojęć, nieuwzględnienie warunku brzegowego, błędna interpretacja wykresu, wybór niewłaściwego wzoru albo zbyt szerokie uogólnienie wyniku.

8. Wróć do materiału po czasie

Jedna poprawna odpowiedź bezpośrednio po otrzymaniu informacji zwrotnej nie dowodzi trwałego opanowania. Wróć do pytania po jednym lub kilku dniach (strategia spaced repetition), ale nie proś o jego dokładne powtórzenie. Nowe zadanie powinno sprawdzać tę samą wiedzę w zmienionym kontekście.

Możesz poprosić model o zaplanowanie zestawów powtórkowych, ale samodzielnie kontroluj harmonogram i zakres materiału. Pytania, które rozwiązujesz poprawnie i pewnie, mogą pojawiać się rzadziej, pytania ujawniające błędy powinny wracać częściej i w różnych formach.

Dostosuj trudność, ale nie usuwaj wysiłku

Stopniowanie wsparcia

Jeżeli nie potrafisz odpowiedzieć, nie prosz od razu o rozwiązanie.

- 1 ponów samodzielną próbę
- 2 poproś o wskazanie luki lub rodzaju błędu
- 3 poproś o jedno pytanie naprowadzające
- 4 poproś o przypomnienie właściwego pojęcia
- 5 dopiero na końcu przeanalizuj pełne rozwiązanie

Po skorzystaniu ze wskazówki wykonaj podobne zadanie bez wsparcia. Prawidłowa odpowiedź z podpowiedzią nie jest jeszcze dowodem samodzielnego opanowania umiejętności.

Prowadź rejestr błędów

Po sesji zapisz:

- czego dotyczyło pytanie
- jakiej odpowiedzi udzieliłeś
- na czym polegał błąd
- dłaczego go popełniłeś
- jaka zasada pozwala go skorygować
- kiedy ponownie sprawdzisz ten obszar

Nie zapisuj wyłącznie poprawnej odpowiedzi. Nazwij mechanizm błędu, na przykład: pomylenie dwóch pojęć, nieuwzględnienie warunku brzegowego, błędna interpretacja wykresu, wybór niewłaściwego wzoru albo zbyt szerokie uogólnienie wyniku.

Wróć do materiału po czasie

spaced repetition

Jedna poprawna odpowiedź bezpośrednio po otrzymaniu informacji zwrotnej nie dowodzi trwałego opanowania. Wróć do pytania po jednym lub kilku dniach, ale nie prosz o jego dokładne powtórzenie. Nowe zadanie powinno sprawdzać tę samą wiedzę w zmienionym kontekście.

poprawnie i pewnie

rzadziej

pytania ujawniające błędy

częściej i w różnych formach

Możesz poprosić system o zaplanowanie zestawów powtórkowych, ale samodzielnie kontroluj harmonogram i zakres materiału.

Prompt: sesja samotestowania

Na podstawie załączonych efektów uczenia się, kryteriów oceniania i materiałów przeprowadź ze mną sesję samotestowania z zakresu: [wpisz zakres].

Korzystaj wyłącznie z dostarczonych materiałów. Jeżeli nie zawierają informacji potrzebnej do oceny odpowiedzi, poinformuj mnie o tym.

Przygotuj pytania sprawdzające odtwarzanie wiedzy, rozumienie, zastosowanie i analizę. Uwzględnij nowe przykłady, wykrywanie błędów, porównywanie rozwiązań i uzasadnianie decyzji.

Zadawaj jedno pytanie naraz. Nie pokazuj odpowiedzi ani podpowiedzi przed moją próbą. Po mojej odpowiedzi poproś mnie o ocenę pewności w skali od 0 do 100%. Następnie wskaż, co jest poprawne, czego brakuje, na czym polega ewentualny błąd i gdzie w materiale mogę to sprawdzić.

Poproś mnie o poprawienie odpowiedzi własnymi słowami. Następnie zadaj nowe pytanie sprawdzające ten sam mechanizm w innym kontekście. Nie zwiększaj poziomu wsparcia, dopóki nie podejmę samodzielnej próby.

Po zakończeniu przygotuj listę zagadnień opanowanych, niepewnych i wymagających ponownej pracy. Nie uznawaj umiejętności za opanowaną wyłącznie na podstawie odpowiedzi udzielonej po podpowiedzi.

OSIEM ZASAD

Samotestowanie z GenAI

- 1 Skuteczne samotestowanie zaczyna się od samodzielnego wydobywania wiedzy z pamięci, bez materiałów, odpowiedzi i podpowiedzi.
- 2 Pytania powinny sprawdzać nie tylko znajomość pojęć, lecz także rozumienie, zastosowanie, analizę błędów i uzasadnianie decyzji.
- 3 Ocena pewności odpowiedzi pomaga wykrywać zarówno luki w wiedzy, jak i błędne przekonania, których nie jesteśmy świadomi.
- 4 Informacja zwrotna powinna wyjaśniać błąd, wskazywać jego źródło i prowadzić do poprawienia odpowiedzi własnymi słowami.
- 5 Odpowiedzi i klucze wygenerowane przez model wymagają sprawdzenia w materiałach obowiązujących na przedmiocie.
- 6 O opanowaniu wiedzy świadczy wykonanie nowego zadania bez podpowiedzi, a nie rozpoznanie odpowiedzi ani zrozumienie cudzego rozwiązania.
- 7 Powrót do materiału po pewnym czasie pozwala sprawdzić trwałość wiedzy i odróżnić chwilową poprawę od rzeczywistego uczenia się.
- 8 GenAI może organizować praktykę i generować kolejne zadania, ale nie powinna zastępować samodzielnego przypominania, rozumowania i podejmowania decyzji.

Dwa tryby pracy

1. Tryb uczenia się

W tym trybie otrzymujesz informację zwrotną po każdym pytaniu, poprawiasz odpowiedź i rozwiązujesz podobny problem. Celem jest wykrywanie i korygowanie błędów.

2. Tryb próby egzaminacyjnej

W tym trybie:

- pracujesz bez materiałów i narzędzi,
- obowiązuje określony czas,
- nie otrzymujesz podpowiedzi ani informacji zwrotnej podczas pracy,
- zadania dotyczą nowych przykładów,
- pytania są wymieszane i nie wskazują, jakiej metody należy użyć,
- odpowiedzi wymagają uzasadnienia,
- ocenę i analizę błędów wykonujesz dopiero po zakończeniu całej próby.

Warunki powinny odpowiadać zasadom rzeczywistego egzaminu. Jeżeli na egzaminie wolno korzystać z tablic, notatek, programu lub kalkulatora, uwzględnij je również podczas próby.

Kiedy możesz uznać, że materiał jest opanowany?

Nie wtedy, gdy odpowiedź wydaje Ci się znajoma ani wtedy, gdy rozumiesz rozwiązanie przedstawione przez model. Mocniejszym dowodem jest to, że po pewnym czasie potrafisz:

- samodzielnie odtworzyć potrzebną wiedzę,
- rozwiązać nowe zadanie bez podpowiedzi,
- dobrać metodę bez wskazania jej nazwy,
- uzasadnić kolejne kroki,
- wykryć błąd w cudzym rozwiązaniu,
- określić granice zastosowania rozwiązania.

Dwa tryby pracy

TRYB PIERWSZY

Tryb uczenia się

Otrzymujesz informację zwrotną po każdym pytaniu, poprawiasz odpowiedź i rozwiązujesz podobny problem.

Celem jest wykrywanie i korygowanie błędów.

TRYB DRUGI

Tryb próby egzaminacyjnej

pracujesz bez materiałów i narzędzi

obowiązuje określony czas

nie otrzymujesz podpowiedzi ani informacji zwrotnej podczas pracy

zadania dotyczą nowych przykładów

pytania są wymieszane i nie wskazują, jakiej metody należy użyć

odpowiedzi wymagają uzasadnienia

ocenę i analizę błędów wykonujesz po zakończeniu całej próby

Warunki powinny odpowiadać zasadom rzeczywistego egzaminu. Jeżeli wolno korzystać z tablic, notatek, programu lub kalkulatora, uwzględnij je również podczas próby.

Kiedy możesz uznać, że materiał jest opanowany?

Nie wtedy, gdy odpowiedź wydaje Ci się znajoma ani wtedy, gdy rozumiesz rozwiązanie przedstawione przez model. Mocniejszym dowodem jest to, że po pewnym czasie potrafisz:

✓ samodzielnie odtworzyć potrzebną wiedzę

✓ rozwiązać nowe zadanie bez podpowiedzi

✓ dobrać metodę bez wskazania jej nazwy

✓ uzasadnić kolejne kroki

✓ wykryć błąd w cudzym rozwiązaniu

✓ określić granice zastosowania rozwiązania

EDU × AI

Podsumowując: najpierw odpowiedz z pamięci, potem sprawdź i popraw, a po czasie rozwiąż nowy problem bez wsparcia. GenAI może organizować tę praktykę, ale nie powinna wykonywać jej za Ciebie.

Podsumowanie modułu 6

Moduł 6 OSIEM WNIOSKÓW

Podsumowanie

1 Kompetencje AI obejmują cztery obszary:

Rozumienie
technologii

Celowe wykorzystanie

Krytyczna ocena
rezultatów

Konsekwencje etyczne
i społeczne

— Zanim zaczniesz

2 Sposób wykorzystania GenAI powinien zależeć od celu zadania i zasad obowiązujących na danym przedmiocie.

3 Warto odróżniać trzy rodzaje czynności:

- wykonywane samodzielnie
- wspierane przez GenAI
- bezpiecznie delegowane

— Przebieg pracy

4 Czytanie, pisanie, rozwiązywanie zadań, programowanie i analiza danych powinny rozpoczynać się od własnego rozpoznania problemu.

5 Wygenerowana odpowiedź wymaga sprawdzenia w źródłach, za pomocą obliczeń, testów lub samodzielnego rozumowania.

6 Praca powinna kończyć się samodzielnym wyjaśnieniem, rozwiązaniem nowego problemu albo zastosowaniem wiedzy bez dostępu do narzędzia.

7 Dziennik wykorzystania GenAI

Ułatwia ocenę własnego procesu oraz przygotowanie przejrzystego oświadczenia.

8 Podpisana praca to Twoja praca

Student pozostaje odpowiedzialny za treść, źródła, kod, interpretację, decyzje i wnioski.

Moduł 7. Kompetencje w epoce post-AI

Czas realizacji: ok. 25 minut

W tym module określisz, które kompetencje warto rozwijać w świecie z AI, które czynności można delegować technologii oraz jak zachować kontrolę nad celem, decyzjami i odpowiedzialnością.

Pytanie otwierające

Którą czynność wykonywaną obecnie samodzielnie możesz bezpiecznie przekazać GenAI, a której nie chcesz delegować, ponieważ jej wykonywanie rozwija kompetencję ważną w Twojej dziedzinie?

— PYTANIE OTWIERAJĄCE



Którą czynność wykonywaną obecnie samodzielnie możesz bezpiecznie przekazać GenAI, a której nie chcesz delegować, ponieważ jej wykonywanie rozwija kompetencję ważną w Twojej dziedzinie?

7.1. Mapa kompetencji

AI zmienia sposób wykonywania zadań, ale nie usuwa potrzeby wiedzy dziedzinowej, rozumowania, oceny jakości, samoregulacji, współpracy i odpowiedzialności. Przeciwnie, im więcej czynności można delegować technologii, tym ważniejsza staje się zdolność rozpoznania, które z nich warto wykonać samodzielnie.

Jedną z podstawowych kompetencji staje się **świadome zarządzanie podziałem pracy między człowiekiem a modelem**: określenie celu, wybór zakresu wsparcia, kontrola procesu, ocena wyniku oraz zachowanie odpowiedzialności za decyzję.

UNESCO ujmuje kompetencje związane z AI jako połączenie podejścia skoncentrowanego na **człowieku, etyki, rozumienia technologii** i jej zastosowań. W ramie dla nauczycieli uwzględniono również kompetencje pedagogiczne i rozwój zawodowy ([Miao i in. 2024](#), [Miao i Cukurova 2024](#)).

Ramy UNESCO zostały opracowane dla uczniów i nauczycieli, dlatego w tym kursie traktujemy je jako punkt odniesienia dostosowany do szkolnictwa wyższego.

OBSZAR KOMPETENCJI	CO OBEJMUJE
Wiedza dziedzinowa	pojęcia, modele, metody i standardy właściwe dla dyscypliny
Rozumowanie i rozwiązywanie problemów	formułowanie problemu, analiza założeń, wnioskowanie i wybór rozwiązania
Ocena jakości	sprawdzanie informacji, źródeł, danych, obliczeń, kodu i argumentów
Samoregulacja i metapoznanie	planowanie, monitorowanie rozumienia, ocena własnego wykonania i zmiana strategii
Kompetencje związane z AI	rozumienie zasad działania, możliwości, ograniczeń i ryzyka systemów
Kompetencje informacyjne i cyfrowe	wyszukiwanie, porównywanie, dokumentowanie i bezpieczne przetwarzanie informacji
Współpraca i komunikacja	uzgadnianie znaczeń, argumentowanie, udzielanie informacji zwrotnej i wspólne podejmowanie decyzji
Odpowiedzialność	rozpoznawanie konsekwencji, ochrona danych, przejrzystość i przyjmowanie odpowiedzialności za decyzje

Ćwiczenie

Wybierz z listy:

- dwie kompetencje, które są szczególnie ważne w Twojej dziedzinie,
- jedną kompetencję, którą technologia może wspierać,
- jedną kompetencję, której nie chcesz osłabić przez nadmierne delegowanie,
- jedną kompetencję, którą chcesz świadomie rozwijać.

Krótko uzasadnij każdy wybór, zapisz swoje przemyślenia.

7.2. Czy tę czynność warto delegować?

Nie każda czynność, którą można wykonać z wykorzystaniem GenAI, powinna zostać technologii powierzona. Decyzja zależy od celu, poziomu wiedzy, ryzyka błędu oraz tego, czy wykonanie danej czynności samo w sobie jest ważnym elementem uczenia się lub pracy zawodowej.

PIĘĆ PYTAŃ

Czy tę czynność warto delegować?

1 Czy wykonanie tej czynności jest celem uczenia się?

Jeśli tak, jej delegowanie może usunąć potrzebną pracę poznawczą.

2 Czy mam wiedzę potrzebną do oceny wyniku?

Jeśli nie, istnieje ryzyko przyjęcia błędnej lub niepełnej odpowiedzi.

3 Jakie mogą być konsekwencje błędu?

Im wyższe ryzyko dla ludzi, środowiska, bezpieczeństwa, prawa lub jakości decyzji, tym większa potrzeba kontroli człowieka.

4 Czy sposób wykorzystania GenAI można przejrzeć i udokumentować?

Brak możliwości ujawnienia udziału technologii jest sygnałem, że jej zastosowanie może być niewłaściwe.

5 Czy po zakończeniu potrafię wyjaśnić rezultat i wykonać podobne zadanie samodzielnie?

Jeśli nie, rezultat nie potwierdza jeszcze posiadania kompetencji.

EDU x AI

Przed delegowaniem zadaj pięć pytań:

1. **Czy wykonanie tej czynności jest celem uczenia się?**
Jeśli tak, jej delegowanie może usunąć potrzebną pracę poznawczą.
2. **Czy mam wiedzę potrzebną do oceny wyniku?**
Jeśli nie, istnieje ryzyko przyjęcia błędnej lub niepełnej odpowiedzi.
3. **Jakie mogą być konsekwencje błędu?**
Im wyższe ryzyko dla ludzi, środowiska, bezpieczeństwa, prawa lub jakości decyzji, tym większa potrzeba kontroli człowieka.
4. **Czy sposób wykorzystania GenAI można przejrzeć i udokumentować?**
Brak możliwości ujawnienia udziału technologii jest sygnałem, że jej zastosowanie może być niewłaściwe.
5. **Czy po zakończeniu potrafię wyjaśnić rezultat i wykonać podobne zadanie samodzielnie?**
Jeśli nie, rezultat nie potwierdza jeszcze posiadania kompetencji.

Trzy możliwe decyzje

DECYZJA	KIEDY JĄ PODJĄĆ	PRZYKŁAD
Wykonuję samodzielnie	czynność jest celem uczenia się albo podstawą późniejszej oceny	sformułowanie hipotezy, przeprowadzenie kluczowego rozumowania
Korzystam ze wsparcia	GenAI dostarcza wskazówek, przykładów lub informacji zwrotnej, ale decyzje pozostają po stronie użytkownika	porównanie dwóch rozwiązań, wskazanie możliwych luk w argumentacji
Deleguję pod kontrolą	czynność jest rutynowa, jej wykonanie nie stanowi celu uczenia się, a rezultat można sprawdzić	zmiana formatu danych, przygotowanie pierwszej wersji tabeli lub transkrypcji

EDU × AI

Ćwiczenie

Granica delegowania

Wybierz jedną czynność, którą często wykonujesz podczas studiowania, prowadzenia zajęć albo pracy naukowej. Zastosuj pięć pytań i zdecyduj:

1. wykonuję ją samodzielnie,
2. korzystam ze wsparcia GenAI,
3. deleguję ją technologii pod kontrolą.

Zapisz najważniejszy argument uzasadniający decyzję.

7.3. Autodiagnoza kompetencji AI

KARTA PRACY

Autodiagnoza kompetencji AI

Oceń każde stwierdzenie w skali:

1 jeszcze nie

2 częściowo

3 zazwyczaj

4 konsekwentnie

Rozumienie technologii

- ☐ potrafię wyjaśnić, w jaki sposób model językowy generuje odpowiedź
- ☐ rozumiem, dlaczego ten sam model może wygenerować różne odpowiedzi
- ☐ znam ograniczenia wykorzystywanego narzędzia
- ☐ rozpoznaję antropomorfizowanie technologii

Uczenie się

- ☐ potrafię określić cel przed użyciem GenAI
- ☐ podejmuję własną próbę przed poproszeniem o rozwiązanie
- ☐ proszę o wskazówki i informację zwrotną
- ☐ sprawdzam wiedzę bez dostępu do narzędzia

Ocena informacji

- ☐ sprawdzam, czy podane publikacje istnieją
- ☐ docieram do źródeł pierwotnych
- ☐ odróżniam wynik badania od jego interpretacji
- ☐ zauważam odpowiedzi potwierdzające założenia zawarte w pytaniu

Odpowiedzialność

- ☐ znam zasady obowiązujące w danym przedmiocie
- ☐ chronię dane osobowe i poufne
- ☐ dokumentuję istotne wykorzystanie GenAI
- ☐ potrafię obronić każdy element podpisanej pracy

Konsekwencje społeczne

- ☐ rozpoznaję stereotypy w wygenerowanych treściach
- ☐ oceniam dostępność i sprawiedliwość wykorzystania narzędzia
- ☐ uwzględniam koszt środowiskowy
- ☐ nie traktuję systemu jako zamiennika relacji i profesjonalnego wsparcia

Interpretacja

Nie sumuj punktów do jednego wyniku. Sprawdź:

które stwierdzenia otrzymały ocenę 1 lub 2

w jakich sytuacjach najczęściej tracisz kontrolę nad przebiegiem pracy

które kompetencje są potrzebne do oceny rezultatów generowanych w Twojej dziedzinie

czego nie da się wiarygodnie ocenić na podstawie samej samooceny

Audyt

Samo **przekonanie** o posiadaniu kompetencji może być nietrafne, Wybierz **dwa obszary** z najniższą oceną i wykonaj krótki audyt.

Sprawdź, czy potrafisz:

- wskazać cel wykorzystania technologii,
- odtworzyć własny tok rozumowania,
- rozpoznać elementy wygenerowane i samodzielnie opracowane,
- zweryfikować najważniejsze twierdzenia,
- uzasadnić przyjęte decyzje,
- wykonać podobne zadanie bez GenAI.

Porównaj wynik audytu z wcześniejszą samooceną.

Jeżeli oceny się różnią, za bardziej wiarygodny uznaj dowód wynikający z działania.

Zaplanuj konkretne działanie rozwojowe.

7.4. Plan rozwoju jednej kompetencji

Wybierz jedną kompetencję, której rozwój będzie miał największe znaczenie dla Twojego studiowania, nauczania albo pracy naukowej. Zamiast ogólnego celu zaplanuj obserwowalne działanie.

1. Kompetencja, którą rozwijam:
2. Sytuacja, w której jej potrzebuję:
3. Działanie, które będę wykonywać:
4. Sposób wykorzystania GenAI:
5. Czynność, której nie będę delegować:
6. Dowód rozwoju kompetencji:
7. Termin sprawdzenia postępu:

Dobrym dowodem rozwoju może być samodzielnie rozwiązane zadanie, uzasadnienie decyzji, poprawiona wersja pracy, analiza błędu, porównanie źródeł albo udokumentowany przebieg działania (samo częstsze korzystanie z narzędzia nie jest dowodem rozwoju kompetencji).

KARTA PRACY

Plan rozwoju jednej kompetencji

Wybierz jedną kompetencję, której rozwój będzie miał największe znaczenie dla Twojego studiowania, nauczania albo pracy naukowej. Zamiast ogólnego celu zaplanuj obserwowalne działanie.

1 Kompetencja, którą rozwijam

2 Sytuacja, w której jej potrzebuję

3 Działanie, które będę wykonywać

4 Sposób wykorzystania GenAI

5 Czynność, której nie będę delegować

6 Dowód rozwoju kompetencji

7 Termin sprawdzenia postępu

Podsumowanie modułu 7

Moduł 7 PIĘĆ WNIOSKÓW

Podsumowanie

- 1** Kompetencje potrzebne w świecie z AI to te same kompetencje, które sprzyjają uczeniu się bez technologii, tylko stosowane w nowych warunkach.

Wiedza dziedzinowa

Rozumowanie

Ocena jakości

Samoregulacja

Współpraca

Odpowiedzialność

Rozpoznawanie granic własnej wiedzy

Komunikacja

- 2** Granice własnej wiedzy

Ekspertyza adaptacyjna i rozpoznawanie granic własnej wiedzy zyskują na znaczeniu, gdy GenAI potrafi wygenerować odpowiedź nawet przy niewystarczających danych.

- 3** Osąd i trafne pytania

Osąd jakości i umiejętność zadawania trafnych pytań decydują o tym, czy praca z GenAI wspiera, czy ogranicza rozwój kompetencji.

- 4** Kiedy delegowanie jest uzasadnione

Im mniejsze znaczenie czynności dla uczenia się i im łatwiej zweryfikować rezultat, tym bardziej uzasadnione jest przekazanie jej technologii.

● kluczowe dla uczenia się, rób sam

● pomocnicze, wspieraj się GenAI

○ łatwe do sprawdzenia, deleguj

- 5** Najważniejsza kompetencja to świadoma decyzja, z czego skorzystać.

KIEDY
korzystać z GenAI

KIEDY
pracować samodzielnie

KIEDY
zwrócić się do człowieka

Słownik pojęć

Słownik pojęć

Poniższy słownik obejmuje terminy techniczne i metodyczne występujące w kursie, uporządkowane alfabetycznie.

AIAS (Artificial Intelligence Assessment Scale, Skala AIAS)

Rama porządkująca dopuszczalny zakres wykorzystania GenAI w zadaniu, obejmująca pięć poziomów: od pracy całkowicie bez GenAI po twórcze badanie możliwości technologii. Służy do projektowania zadań i komunikowania zasad studentom, nie do wykrywania użycia GenAI. aiassessmentscale.com

Antropomorfizacja

Przypisywanie technologii cech właściwych ludziom (intencji, emocji, troski, świadomości) na podstawie płynnej, konwersacyjnej formy odpowiedzi. Model nie doświadcza emocji ani nie rozumie sytuacji użytkownika; wypowiedzi przypominające empatię są efektem wzorców językowych, a nie stanu wewnętrznego systemu.

Chunking (grupowanie informacji)

Strategia poznawcza polegająca na łączeniu pojedynczych informacji w większe jednostki znaczeniowe, co zwiększa efektywną pojemność pamięci roboczej. Osoby o rozwiniętej wiedzy dziedzinowej grupują informacje sprawniej niż osoby początkujące.

Cognitive offloading (odciążenie poznawcze)

Przeniesienie części pracy poznawczej na zewnętrzne narzędzie w celu zmniejszenia wysiłku potrzebnego do wykonania zadania (np. kalkulator, nawigacja, GenAI). Nie jest z definicji niekorzystne, problem pojawia się, gdy narzędzie przejmuje czynność potrzebną do osiągnięcia efektu uczenia się.

Dialog sokratejski

Sposób prowadzenia rozmowy oparty na pytaniach prowadzących rozmówcę do samodzielnego zauważenia założenia, sprzeczności lub brakującego dowodu, zamiast bezpośredniego podania odpowiedzi. W kontekście GenAI oznacza zaprojektowanie interakcji tak, by model zadawał pytania, a nie generował rozwiązania.

Duży model językowy (LLM, Large Language Model)

Model wytrenowany na bardzo dużych zbiorach tekstu, który podczas generowania oblicza prawdopodobieństwo kolejnych tokenów i wybiera je zgodnie z kontekstem oraz ustawieniami systemu. LLM jest elementem systemu działającego za interfejsem chatbota, nie jest z nim tożsamy.

Ekspertyza adaptacyjna oraz ekspertyza rutynowa

Ekspertyza rutynowa pozwala sprawnie wykonywać znane zadania według wyuczonej procedury. Ekspertyza adaptacyjna pozwala modyfikować metody i tworzyć nowe rozwiązania w zmieniających się warunkach, jest szczególnie istotna, gdy GenAI generuje znane schematy, które nie pasują do nowego kontekstu.

GenAI (generatywna sztuczna inteligencja)

Podzbiór sztucznej inteligencji obejmujący modele służące do generowania nowych treści (tekstu, kodu, obrazów, dźwięku, wideo) na podstawie wzorców wykrytych w danych treningowych oraz informacji przekazanych przez użytkownika. Nie każdy system AI jest generatywny.

Iluzja kompetencji (iluzja mistrzostwa)

Błędne przekonanie, że łatwość czytania, rozpoznawania lub śledzenia gotowego rozwiązania jest dowodem opanowania wiedzy. Rozpoznanie poprawnie brzmiącej odpowiedzi jest łatwiejsze niż jej samodzielne odtworzenie lub zastosowanie, co może prowadzić do zawyżonej samooceny uczenia się.

Kompetencje AI (AI literacy)

Zestaw umiejętności pozwalających rozumieć podstawy działania AI, celowo wykorzystywać technologię, krytycznie oceniać jej rezultaty oraz rozpoznawać konsekwencje etyczne i społeczne jej stosowania. Obejmuje coś więcej niż sprawne posługiwanie się narzędziem.

Lenistwo metapoznawcze (metacognitive laziness)

Możliwe zjawisko ograniczenia działań metapoznawczych (planowania, monitorowania i oceny własnej pracy), gdy część tych czynności zostaje przekazana technologii. Nie jest cechą trwałą osoby ani diagnozą psychologiczną, lecz opisem sposobu pracy w określonych warunkach; wymaga ostrożnego stosowania.

Neuroplastyczność

Zdolność mózgu do zmiany siły i liczby połączeń między neuronami w odpowiedzi na doświadczenie i aktywność poznawczą, utrzymująca się przez całe życie. Sposób pracy z technologią, a nie tylko dostęp do niej, wpływa na to, jakie ślady pamięciowe powstają.

Osąd ewaluacyjny (evaluative judgement)

Zdolność podejmowania trafnych decyzji dotyczących jakości własnej pracy i pracy innych osób, w tym treści wygenerowanych przez model. Rozwija się przez podejmowanie decyzji i obserwowanie ich konsekwencji, nie tylko przez poznawanie kryteriów.

Pamięć długotrwała

System pamięci bez praktycznych ograniczeń pojemności, w którym informacje mogą być przechowywane przez całe życie. Celem uczenia się jest przeniesienie informacji z pamięci roboczej do pamięci długotrwałej poprzez aktywne przetwarzanie, przywoływanie i powtórki rozłożone w czasie.

Pamięć robocza

System pozwalający przez krótki czas przechowywać i przetwarzać informacje potrzebne do bieżącego działania; ma ograniczoną pojemność, którą można efektywnie zwiększyć przez chunking.

Paradoks pamięci

Zjawisko opisujące sytuację, w której im bardziej zaawansowane stają się zewnętrzne narzędzia informacyjne, tym mniej potrzebne wydaje się budowanie własnej wiedzy, mimo że korzystanie z tych narzędzi wymaga wiedzy własnej do formułowania pytań, oceny odpowiedzi i wykrywania błędów.

Pochlebczość modeli (sykofancja, sycophancy)

Tendencja modelu do generowania odpowiedzi zgodnych z przekonaniami, opiniami lub oczekiwaniami użytkownika, nawet gdy prowadzi to do treści niezgodnej z faktami lub słabiej uzasadnionej. Może przyjmować formę potwierdzania założeń zawartych w pytaniu lub zmiany odpowiedzi po sprzeciwie użytkownika, bez pojawienia się nowych danych.

Retrieval practice (przywoływanie z pamięci, autotestowanie)

Technika uczenia się polegająca na samodzielnym przypominaniu sobie informacji z pamięci, np. w formie quizu lub próby odtworzenia treści bez podglądania materiału. Należy do technik wysokiej użyteczności według przeglądu, poprawia trwałość zapamiętywania skuteczniej niż wielokrotne czytanie.

Samoregulacja uczenia się (Self-Regulated Learning, SRL)

Aktywne kierowanie własnymi procesami poznawczymi, motywacyjnymi, emocjonalnymi i behawioralnymi w celu osiągnięcia efektu uczenia się, obejmujące planowanie, monitorowanie postępów i ocenę rezultatu. GenAI może ją wspierać (pytania, propozycje strategii) lub ograniczać (przejęcie planowania i oceny).

Spaced repetition (powtórki rozłożone w czasie)

Technika uczenia się polegająca na powtarzaniu materiału w coraz większych odstępach czasu. Wraz z retrieval practice należy do technik wysokiej użyteczności.

Stopień ingerencji GenAI (niski lub wysoki)

Klasyfikacja przyjęta w wytycznych PG, rozróżniająca zastosowania GenAI związane głównie z poprawianiem i przetwarzaniem materiału przygotowanego przez człowieka (niski stopień ingerencji) od zastosowań, w których model generuje elementy stanowiące istotną część rezultatu, np. argumenty, kod czy interpretację wyników (wysoki stopień ingerencji). Determinuje rodzaj wymaganego oświadczenia.

Token

Podstawowa jednostka przetwarzana przez model językowy podczas generowania tekstu, może nią być całe słowo, jego część, znak interpunkcyjny lub inny element. Model generuje tekst sekwencyjnie, token po tokenie.

Transformer

Architektura sieci neuronowej wykorzystująca mechanizm uwagi (attention) do przetwarzania zależności między elementami sekwencji tekstu; stanowi podstawę większości współczesnych dużych modeli językowych.

Uczenie maszynowe

Metody, dzięki którym modele wykrywają wzorce w danych treningowych i wykorzystują je do wykonywania określonych zadań, bez konieczności ręcznego zapisania wszystkich reguł działania przez programistę. GenAI jest jednym z zastosowań uczenia maszynowego.

Bibliografia

- Abercrombie, G., Cercas Curry, A., Dinkar, T., Rieser, V. & Talat, Z. (2023). Mirages: On anthropomorphism in dialogue systems. W Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing (s. 4776-4790). Association for Computational Linguistics.
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakcı, Ö. & Mariman, R. (2025). Generative AI without guardrails can harm learning: Evidence from high school mathematics. Proceedings of the National Academy of Sciences, 122(26), e2422633122.
- Bender, E. M., Gebru, T., McMillan-Major, A. & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? W Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (s. 610-623). Association for Computing Machinery.
- Brady, W. J., Wills, J. A., Jost, J. T., Tucker, J. A. & Van Bavel, J. J. (2017). Emotion shapes the diffusion of moralized content in social networks. Proceedings of the National Academy of Sciences, 114(28), 7313-7318.
- Chatfield, T. (2025). AI and the future of pedagogy [[White paper](#)]. Sage.
- Committee on Publication Ethics. (2023). Authorship and AI tools. COPE position statement. <https://publicationethics.org/cope-position-statements/ai-author>
- Cowan, N. (2001). The magical number 4 in short-term memory: A reconsideration of mental storage capacity. Behavioral and Brain Sciences, 24(1), 87-114.
- De Freitas, J., Uguralp, A. K., Uguralp, Z. O. & Puntoni, S. (2024). AI companions reduce loneliness. Journal of Consumer Research, 52(6), 1126-1148.
- Deshpande, A., Rajpurohit, T., Narasimhan, K. & Kalyan, A. (2023). Anthropomorphization of AI: Opportunities and risks. arXiv:2305.14784. <https://doi.org/10.48550/arXiv.2305.14784>
- Deslauriers, L., McCarty, L. S., Miller, K., Callaghan, K. & Kestin, G. (2019). Measuring actual learning versus feeling of learning in response to being actively engaged in the classroom. Proceedings of the National Academy of Sciences, 116(39), 19251-19257.
- Etkin, H. K., Etkin, K. J., Carter, R. J. & Rolle, C. E. (2025). Differential effects of GPT-based tools on comprehension of standardized passages. Frontiers in Education, 10, Article 1506752.
- Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y., Shen, Y., Li, X. & Gašević, D. (2024). Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. British Journal of Educational Technology, 56(2), 489-530.
- Gimpel, H., Hall, K., Decker, S., Eymann, T., Lämmermann, L., Mädche, A., Röglinger, M., Ruiner, C., Schoch, M., Schoop, M., Urbach, N., & Vandrik, S. (2023). Unlocking the power of generative AI models and systems such as GPT-4 and ChatGPT for higher education: A guide for students and lecturers. Hohenheim Discussion Papers in Business, Economics and Social Sciences, 02-2023. University of Hohenheim.

Giray, L. (2026). When using AI in scientific research: Start with human, end with human. *TechTrends*, 70(1), 265-272.

Illingworth, S., & Forsyth, R. (2026). *GenAI in higher education: Redefining teaching and learning*. Bloomsbury Academic.

International Committee of Medical Journal Editors. (2024). Recommendations for the conduct, reporting, editing, and publication of scholarly work in medical journals. <https://www.icmje.org/recommendations/>

Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A. & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248.

Lai, L., Pan, Y., Xu, R. & Jiang, Y. (2025). Depression and the use of conversational AI for companionship among college students: The mediating role of loneliness and the moderating effects of gender and mind perception. *Frontiers in Public Health*, 13, Article 1580826.

Li, P., Yang, J., Islam, M. A. & Ren, S. (2025). Making AI less “thirsty”: Uncovering and addressing the secret water footprint of AI models. *Communications of the ACM*, 68(7), 54-61

Liang, W., Yuksekgonul, M., Mao, Y., Wu, E. & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), Article 100779.

Long, D. & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. W *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (s. 1-16). Association for Computing Machinery.

Luccioni, A. S., Jernite, Y. & Strubell, E. (2024). Power hungry processing: Watts driving the cost of AI deployment? W *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency* (s. 85-99). Association for Computing Machinery.

Miao, F. & Cukurova, M. (2024). AI competency framework for teachers. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000391104>

Miao, F. & Holmes, W. (2023). Guidance for generative AI in education and research. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000386693>

Miao, F. & Shiohira, K. (2024). AI competency framework for students. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000391105>

Mollick, E. R. & Mollick, L. (2023). Assigning AI: Seven Approaches for Students, with Prompts (September 23, 2023). The Wharton School Research Paper, Available at SSRN: <https://ssrn.com/abstract=4475995> or <http://dx.doi.org/10.2139/ssrn.4475995>

Ng, D. T. K., Leung, J. K. L., Chu, S. K. W. & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041.

Oakley, B., Johnston, M., Chen, K.-Z., Jung, E., & Sejnowski, T. J. (2025). The memory paradox: Why our brains need knowledge in an age of AI. In M. Rangeley & N. Fairfax

(Eds.), *The artificial intelligence revolution: Challenges and opportunities* (pp. 573–628). Springer.

Peters, U. & Chin-Yee, B. (2025). Generalization bias in large language model summarization of scientific research. *Royal Society Open Science*, 12(4), Article 241776.

Politechnika Gdańska. (2024). Wytyczne dotyczące stosowania narzędzi generatywnej sztucznej inteligencji na Politechnice Gdańskiej. Załącznik do Pisma Okólnego Rektora Politechniki Gdańskiej nr 29/2024.

Richardson, I. (2025, November 11). Universities risk irrelevance by failing to engage fully with AI. [Times Higher Education](https://www.timeshighereducation.com/news/universities-risk-irrelevance-by-failing-to-engage-fully-with-ai).

Risko, E. F. & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676-688.

Scholich, T., Barr, M., Wiltsey Stirman, S., & Raj, S. (2025). A comparison of responses from human therapists and large language model-based chatbots to mental health scenarios. *JMIR Mental Health*, 12, Article e69709.

Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Aspell, A., Bowman, S. R., Cheng, N., Durmus, E., Hatfield-Dodds, Z., Johnston, S. R., Kravec, S., Maxwell, T., McCandlish, S., Ndousse, K., Rausch, O., Schiefer, N., Yan, D., Zhang, M., & Perez, E. (2023). Towards understanding sycophancy in language models. In *The Twelfth International Conference on Learning Representations*. <https://arxiv.org/abs/2310.13548>

Soderstrom, N. C. & Bjork, R. A. (2015). Learning versus performance: An integrative review. *Perspectives on Psychological Science*, 10(2), 176-199.

Stromberg, D., Lei, V., & Wu, Y. (2026). The Generative AI Learning Penalty: Evidence from Chinese Secondary Education [preprint] SSRN: <https://ssrn.com/abstract=6868618>

Tang, L., Sun, Z., Ilday, B., Nestor, J. G., Soroush, A., Elias, P. A. & in. (2023). Evaluating large language models on medical evidence summarization. *npj Digital Medicine*, 6, Article 158.

UNESCO. (2021). Recommendation on the ethics of artificial intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>

UNESCO Women for Ethical AI. (2024). Outlook study on artificial intelligence and gender. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000391719>

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł. & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30(pp. 5998-6008). Curran Associates.

Vosoughi, S., Roy, D. & Aral, S. (2018). The spread of true and false news online. *Science*, 359(6380), 1146-1151.

Wang, P., Zhang, L.-Y., Tzachor, A. & Chen, W.-Q. (2024). E-waste challenges of generative artificial intelligence. *Nature Computational Science*, 4(11), 818-823.

Wardle, C. & Derakhshan, H. (2017). Information disorder: Toward an interdisciplinary framework for research and policy making. Council of Europe.

<https://shorensteincenter.org/wp-content/uploads/2017/10/Information-Disorder-Toward-an-interdisciplinary-framework.pdf>

Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P. & Waddington, L. (2023). Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19, Article 26.

Zimmerman, B. J. (2002). Becoming a self-regulated learner: An overview. *Theory Into Practice*, 41(2), 64-70.