



Fundusze Europejskie
dla Rozwoju Społecznego



Rzeczpospolita
Polska

Dofinansowane przez
Unię Europejską



Projekt „Administrowanie UKW w 4K” – FERS.01.05-IP.08-0319/25-00

Uniwersytet Kazimierza Wielkiego w
Bydgoszczy

Transkrypcja szkolenia

Moduł cyfrowy, Szkolenie AI dla administracji – poziom
podstawowy

Informacje projektowe

Materiał opracowany w ramach projektu „Administrowanie UKW w 4K”

nr FERS.01.05-IP.08-0319/25-00

Projekt współfinansowany ze środków Europejskiego Funduszu Społecznego Plus oraz budżetu państwa w ramach programu

Fundusze Europejskie dla Rozwoju Społecznego 2021–2027.

Beneficjent projektu: Uniwersytet Kazimierza Wielkiego w Bydgoszczy

Informacje licencyjne

Autor: Adrianna Piszcz

Licencja: Creative Commons Uznanie autorstwa 4.0 Międzynarodowe (CC BY 4.0)

Materiał może być kopiowany, adaptowany i rozpowszechniany pod warunkiem wskazania autorstwa.

<https://creativecommons.org/licenses/by/4.0/deed.pl>

MODUŁ I: Fundamenty AI w administracji – komunikacja z modelem i procedury bezpieczeństwa

[00:00:02] Wprowadzenie i kwestie formalne

Dzień dobry. Witam serdecznie na szkoleniu AI dla administracji – poziom podstawowy. Ja nazywam się Adrianna Piszcz i będę dla Państwa przewodnikiem w tych pierwszych krokach w pracy z AI. Z kwestii formalnych: materiał szkoleniowy jest dostępny dla Państwa na licencji Creative Commons 4.0. Oznacza to, że można swobodnie wykorzystywać i modyfikować materiał pod warunkiem zachowania informacji o autorstwie.

Nasze szkolenie realizowane jest w ramach projektu „Administrowanie UKW w 4K” i jest częścią modułu cyfrowego. Projekt ten realizowany jest w ramach programu Fundusze Europejskie dla Rozwoju Społecznego na lata 2021–2027, współfinansowanego ze środków Europejskiego Funduszu Społecznego Plus oraz budżetu państwa.

[00:00:56] O autorze szkolenia

Na sam początek chciałabym parę słów o sobie, autorze szkolenia. Pracuję na Wydziale Informatyki w Katedrze Systemów Inteligentnych i Teleinformatyki. Z wykształcenia jestem informatykiem. Oprócz tego mam ukończone studia podyplomowe uprawniające mnie do pracy z dziećmi i młodzieżą. To właśnie tak rozpoczęła się moja przygoda z dydaktyką w roku 2018, a 2 lata później rozpoczęłam pracę na uczelni.

Tutaj dla Państwa chcę połączyć takie dwa filary mojego doświadczenia: filar praktyczny – jestem inżynierem, jestem praktykiem wdrożeniowcem, mam doświadczenie w pracy z kodem i z algorytmami AI; i oczywiście filar akademicki, czyli znajomość specyfiki pracy na uczelni, pewnych procedur, z którymi mam do czynienia. Oprócz tego jestem też certyfikowanym felinoterapeutą – jest to tak naprawdę moja pasja. Dużą wagę przykładam do bezstresowej, przyjaznej atmosfery na zajęciach i szkoleniach, i takim jednym z elementów jest właśnie metodyka pracy edukacyjnej z udziałem kota.

[00:02:12] Struktura szkolenia

Nasze szkolenie będzie się składało z 3 części. Każda część zajmie około 60 minut. Pierwszy moduł to fundamenty AI w administracji, czyli komunikacja z modelem i procedury bezpieczeństwa. Moduł drugi to przetwarzanie wiedzy i tworzenie dokumentacji, a więc praca z własnymi plikami i z interaktywnymi edytorami AI. Moduł trzeci to własny asystent AI z automatyzacją procesów, personalizacja narzędzi i praca z nagraniami oraz syntezą treści.

[00:02:51] Zakres pierwszego modułu

Zacznijmy od tego pierwszego modułu. Tutaj chciałabym Państwu przybliżyć, co konkretnie kryje się w tych podstawach sztucznej inteligencji, w komunikacji z

modelem i w procedurach bezpieczeństwa. Będziemy mówić sobie o tym, czym w ogóle jest ta sztuczna inteligencja, jak ona działa, czym jest prompt. Troszkę będę chciała Państwu odczarować pewne obawy w związku z używaniem sztucznej inteligencji, a więc powiemy o tych mitach, lękach i przede wszystkim bezpiecznym użytkowaniu sztucznej inteligencji – jak to wygląda z kwestiami RODO, o prywatności i anonimowości.

Powiemy sobie o ustawieniach w narzędziach sztucznej inteligencji, pewnych przełącznikach, które będą tutaj dla nas istotne. Będziemy mówili o technikach rozpoznawania halucynacji i weryfikowaniu odpowiedzi. W końcu – czy promptowanie w ogóle się opłaca? Rachunek zysków i strat czasowych. No i wybór narzędzia oraz pierwsza praca na czacie, czyli troszkę praktyki.

[00:04:09] Podstawowe pojęcia: AI, Generatywne AI, LLM

Na sam początek chciałabym przedstawić kilka definicji, czyli tak naprawdę podstawą AI jest zrozumienie pojęć, które będą się tutaj przewijały nie tylko na moim szkoleniu, ale po prostu, jeżeli Państwo zaczną mieć styczność tutaj ze sztuczną inteligencją, to te pojęcia cały czas tam gdzieś będą się przewijać. No i tutaj pierwszym głównym pojęciem jest po prostu sztuczna inteligencja bądź AI od Artificial Intelligence. Mamy takie pojęcie jak generatywna AI. Dostyc często będą się Państwo spotykali z określeniem LLM, czyli duże modele językowe. Nie bez powodu na grafice jest to przedstawione w formie takiego leja, ponieważ te pojęcia się bardzo mocno przeplatają.

[12:30] Na slajdzie przedstawiono grafikę w kształcie leja ilustrującą zależność pojęciową: od najszerzego pojęcia Sztucznej Inteligencji (AI), przez Generatywną AI (GenAI), aż po Duże Modele Językowe (LLM).

Więc czym jest ta sztuczna inteligencja? Jest to szerokie pojęcie, ale ogólnie oznacza zdolność maszyn do wykonywania zadań wymagających ludzkiego myślenia. Nie należy myśleć, że jest to jakiś całkowicie świadomy, odrębny byt. Jest to po prostu zaawansowany program, który jest w stanie rozpoznawać wzorce w bardzo dużych ilościach danych.

Czym jest zatem to generatywne AI? Możecie się spotkać również z takim angielskim odpowiednikiem GenAI. Tradycyjna sztuczna inteligencja po prostu klasyfikuje i analizuje – na przykład jest w stanie nam odfiltrować spam w poczcie mailowej czy wykryć podejrzaną transakcję kartą. Natomiast GenAI od samej nazwy coś generuje, a więc tworzy nam zupełnie nowe treści. Generuje nowe teksty, może generować obrazy, podsumowania, całkowicie nowe treści na podstawie określonych wzorców i wymagań. Tutaj zastosowanie myślę, że jest oczywiste: możemy dzięki temu szybciej zredagować pewne pisma, streszczać duże pliki, duże dokumenty, wyciągać łatwo wnioski czy po prostu wyszukiwać pewne informacje dużo prościej niż poprzez lupkę w PDFie.

Czym są w takim razie te LLM-y, czyli duże modele językowe (Large Language Models)? Jest to model wyszkolony na miliardach tekstów z całego świata, na danych dostępnych w Internecie i jego zadaniem jest przewidywanie kolejnego

najbardziej

prawdopodobnego słowa w kontekście zadanego pytania. Należy na to patrzeć jak na takie bardzo zaawansowane autouzupełnianie w telefonie – kiedy wpisujemy jakiś początek zdania, to telefon nam podpowiada, co może być następne. Działa na podobnej zasadzie, ale po prostu na poziomie całych akapitów, całych konkretnych wiadomości, zachowując styl taki, jaki zadamy w prompcie, jaki został ustalony na wejściu.

[00:07:34] Czym jest prompt i zasada GIGO?

Użyłam słowa prompt, a więc czym jest ten cały prompt? Prompt to jest instrukcja, to co wpisujemy w okienko czatu. Bardzo często na początek wpisujemy prompty zbyt ogólne, które nie przedstawiają dokładnie tego, co chcemy, no i wtedy odpowiedzi również są ogólne, mogą zawierać po prostu błędy. Więc będziemy mówili sobie za chwileczkę o konkretnych krokach, które musimy podjąć, żeby ten prompt był jak najlepszy, żeby dał nam jak najlepszą odpowiedź. Bardzo ogólnikowo mówimy o tym, że ten prompt ma być precyzyjny, a więc musimy określić, jaką rolę ma przyjąć model, jaki styl, jakie mają być dane przekazane w tym piśmie, o co właściwie w ogóle chodzi.

Czemu to jest aż tak ważne? Ten prompting, czyli to co wpisujemy, kieruje się bardzo prostą zasadą: jeżeli wpisujemy byle co, to byle co dostaniemy. Jest to zasada Garbage In, Garbage Out. AI nie jest w stanie domyślić się dokładnie, czego od niego oczekujemy. Jeżeli nie przedstawimy wszystkich informacji, pełnego kontekstu, no to on będzie musiał zgadywać, będzie musiał dopisać informacje, których mu nie przekazaliśmy, a to bardzo często będzie się wiązało z nieprawidłowymi odpowiedziami, nie tym, co faktycznie oczekujemy.

[00:09:19] Cztery elementy skutecznego prompta

Są 4 podstawowe elementy, które powinniśmy zawrzeć zawsze, gdy formułujemy pytanie do sztucznej inteligencji. Pierwszy element to rola – kim ma być AI i kim my jesteśmy w danym zadaniu. Drugi punkt: zadanie – co ma robić, nad czym w tej chwili pracujemy. Trzeci punkt: kontekst – jakie są zasady i tło. Czwarty punkt: format – czy to ma być tekst, tabela, lista, obrazek; jeśli tekst to jak długi i w jakim stylu, czy ma mieć wypunktowania czy myślniki. Wszystkie informacje, które są dla nas istotne, należy zawrzeć.

[12:30] Na slajdzie przedstawiono przykładową strukturę prompta zawiązanego z pracą w dziale kadr uczelni wyższej:

- **Rola:** Jesteś doświadczonym pracownikiem działu kadr na uczelni wyższej.
- **Zadanie:** Napisz projekt odpowiedzi na wniosek pracownika o dofinansowanie z ZFŚS.
- **Kontekst:** Pracownik złożył wniosek po terminie i nie dołączył oświadczenia o dochodach za zeszły rok. Ton pisma powinien być życzliwy, ale stanowczy, wskazując możliwość ponownego złożenia wniosku w kolejnym naborze za 3 miesiące.

- **Format:** Oficjalny projekt pisma, maksymalnie 3 akapity plus wykaz brakujących dokumentów w formie listy punktowej.

[00:10:57] Cztery rodzaje promptów

Mamy 4 rodzaje promptów i tak naprawdę to, jaki rodzaj prompta i styl instrukcji będziemy chcieli wykorzystać, zależy od skomplikowania sprawy. Pierwszy rodzaj, najprostszy, to jest prompt wykonawczy. Drugi rodzaj to prompt rolowy, następnie prompt wzorcowy i prompt krok po kroku (czyli sekwencyjny).

Ten najprostszy, najkrótszy prompt wykonawczy to jest prompt zastosowania przy krótkich, jednorazowych zadaniach. Tutaj mamy coś natychmiast wykonać. Ma to być krótka, bezpośrednia instrukcja, bez konieczności budowania całego kontekstu, tła i znajomości wszystkich szczegółów fabuły naszego zadania. Na przykład: *„Popraw błędy interpunkcyjne w poniższym tekście i skróć go o 20%”*.

Prompt rolowy jest bardzo często wykorzystywany do pism formalnych, decyzji czy korespondencji urzędowej. Tutaj narzucamy już konkretną tożsamość, mamy opisaną swoją sylwetkę, odbiorcę oraz wymagany ton wypowiedzi. Przykład: *„Jesteś doświadczonym pracownikiem kadr. Napisz odpowiedź na wniosek pracownika o dofinansowanie, dbając o formalny, ale życzliwy ton”*. W wielu sytuacjach takie prompty rolowe są bardzo rozbudowane i mogą zająć nawet stronę A4.

Prompt wzorcowy – szczerze mówiąc, dosyć często używam tego rodzaju prompta, jeżeli zależy nam na zachowaniu dokładnie takiego formatu, którego już gdzieś używaliśmy. Jeżeli chcemy zautomatyzować coś, co do tej pory robiliśmy ręcznie, to my już mamy wypracowany jakiś swój schemat i ciężko nam przełamać się i powiedzieć: „dobra, ja tego nie będę robić, oddam to AI”, bo mamy poczucie, że robimy to dobrze. Więc jeśli dobrze prześlemy to sztucznej inteligencji, ona wykona to w bardzo podobny sposób. Podajemy modelowi wzór – poprzednie pismo lub maila – a następnie prosimy o wygenerowanie nowego tekstu według tego samego schematu. Przykład: *„Oto wzór akceptowalnej decyzji odmownej na naszej uczelni. Na jego podstawie przygotuj analogiczne pismo dla sprawy B”*.

Ostatni rodzaj to prompt krok po kroku. Tutaj mówimy o bardzo złożonych problemach, o trudnych przypadkach, kiedy włączamy swoje myślenie i cykl przyczynowo-skutkowy. Opowiadamy, w jaki sposób my byśmy to zrobili i proponujemy, aby AI zrobiła to w taki sam sposób. Dzielimy duże zadanie na etapy i opisujemy każdy etap. Bardzo często taki prompt krok po kroku warto podzielić nawet na kilka różnych wiadomości w czacie. Jak miałam bardzo złożony problem, to wstawienie nawet bardzo długiego opisu w jednej wiadomości nie dawało oczekiwanej odpowiedzi. Ale kiedy wstawiłam najpierw pierwszy krok, a potem napisałam: „zobacz, masz już to co potrzebujesz do kroku drugiego, zrób krok drugi” – i w następnej wiadomości: „zrób krok trzeci” – to odpowiedź finalnie była bardzo satysfakcjonująca. Warto więc prowadzić rozmowę i nadzorować poszczególne etapy pracy.

[00:16:09] Rola AI w pracy administracyjnej

Bardzo ważna kwestia: nie należy się bać utraty pracy przez AI i czystą automatyzację. Należy się bać ludzi, którzy kształcą się z pracy z AI, bo oni będą w stanie te prace wykonywać szybciej i efektywniej. W administracji uczelni mówimy o tym, żeby wyeliminować przede wszystkim zadania powtarzalne, rutynowe, wypracować sobie schemat pracy asystentów i agentów, którzy będą robić pewne rzeczy za nas albo wspomagać naszą pracę. Dzięki temu oszczędzimy trochę czasu i mamy więcej przestrzeni na te obszary, gdzie naprawdę wymagany jest człowiek – empatia i podejmowanie decyzji.

To jest bardzo ważny element w tym szkoleniu: chciałam Państwu uświadomić, że nie ma konieczności gonić za każdym nowym narzędziem AI, które powstaje, ani dopasowywać narzędzia idealnie do danego problemu. Najważniejsze jest, żeby w ogóle zacząć. Nikt nie będzie od Państwa wymagał zakładania 50 kont na różnych platformach. Wystarczy wybrać jedno lub ewentualnie dwa narzędzia, które się Państwu spodobać i dobrze z nimi zaznajomić.

Przedstawiam Państwu analogię do kont w social mediach: ile mają Państwo kont (np. Facebook, Instagram), a z ilu realnie korzystają Państwo na co dzień?

Najczęściej ta liczba jest ograniczona do 1–2 kont. Dlaczego wybór jednego narzędzia jest tak ważny? Daje to możliwość zrozumienia naszego stylu – model uczy się cały czas naszych preferencji, słownictwa i sposobu formułowania myśli. Mamy dzięki temu dużą spójność wyników, przewidywalność, stabilizujemy efekty pracy, obniżamy frustrację i mamy spokój. Jak zbudujemy porządnie jedno silne środowisko robocze, to nie będziemy mieli potrzeby szukać czegoś innego.

[00:19:47] Praktyka i przełamywanie oporów

Jak to wygląda w codziennej pracy? Początki mogą być frustrujące, bo może się okazać, że pisanie właściwej instrukcji zajmie więcej czasu niż ręczne wykonanie zadania. Natomiast jeśli ta praca później się powtarza, to w kolejnych dniach mamy wypracowany konkretny szablon i to samo zadanie wykonamy już w kilkanaście czy kilkadziesiąt sekund. Dużym problemem są nasze przyzwyczajenia i opór.

Początkowo mieliśmy opór przed mailem czy systemami USOS, a w tej chwili mamy opory przed narzędziem, które częściowo ma przejąć naszą pracę. Jest to całkowicie naturalna obawa. AI wymaga coraz mniej umiejętności technicznych, bo rozmawia się z nią jak ze zwykłym człowiekiem na czacie ze znajomym. Przypomina to zatrudnienie asystenta, którego chcemy przyuczyć i odciążyć siebie.

[00:21:44] Prywatność i bezpieczeństwo danych (RODO)

Co się dzieje z rzeczami, które wpisujemy do czatu? Czy ktoś widzi moje pytania? Trzeba mieć świadomość, że czat jest prywatny – żaden przełożony czy znajomy nie zobaczy, co wpisujemy, nikt nie dostaje z tego raportu. Jednak informacje z naszego komputera przesłane są do infrastruktury dostawcy modelu na serwer, gdzie są przetwarzane przez superkomputery. Po co na serwerze? Bo to złożony model dający o wiele większe możliwości niż modele lokalne na małym laptopie.

W większości modeli domyślnie włączona jest opcja wykorzystywania wprowadzanych danych do trenowania i ulepszania modeli. W pracy administracyjnej zalecam wyłączenie tej opcji suwakiem w ustawieniach. Nie chodzi o to, że ktoś z firmy AI czyta nasze pytania, ale o zabezpieczenie prawne – przekazując dane wrażliwe, wizerunek czy pracę twórczą musimy zastanowić się, czy możemy przekazać je usłudze zewnętrznej.

Najważniejszy aspekt prawny RODO to anonimizacja – zasłaniamy lub wycinamy dane wrażliwe. RODO reguluje samo przetwarzanie danych przez system, a nie czytanie ich przez człowieka. Wszystkie imiona, nazwiska, PESEL i adresy zamieniamy na przykładowe dane, np. „Student X” albo „Jan Kowalski”. Gdy pracujemy na PDFie lub zdjęciu, gdzie nie da się wyciąć tekstu, robimy zrzut ekranu przez Narzędzie Wycinanie w Windowsie i zamazujemy dane mazakiem.

[12:30] Na slajdzie przedstawiono podział danych na strefy bezpieczeństwa:

- **Zielona strefa (spokojnie wrzucamy):** To co publiczne w Internecie – uchwały, regulaminy, instrukcje, wytyczne PKA, ogólnodostępne wzory pism, treści ze strony internetowej.
- **Żółta strefa (najpierw zanonimizować):** Projekty pism z konkretnymi danymi osobowymi, maile od interesantów z podpisami, robocze notatki z zebrań z danymi wrażliwymi.
- **Czerwona strefa (absolutnie nie):** Dane medyczne, nieujawnione dane finansowe, tajemnice państwowe i służbowe.

Jeśli uczenie jest włączone, nasz tekst jest wykorzystywany do trenowania kolejnych wersji modelu. Jeśli wyłączymy uczenie, tekst będzie przetwarzany tylko do udzielenia odpowiedzi i skasowany zaraz po jej wysłaniu.

[00:29:11] Pamięć profilowa i zarządzanie nią

Ustawienia to miejsce, gdzie warto zajrzeć na samym początku. Najważniejszym elementem jest włączenie pamięci – chcemy, żeby pewne informacje na nasz temat zostały zapamiętane pomiędzy czatami. Dzięki temu model szybciej rozumie kontekst i udziela odpowiedzi dopasowanych do naszego stylu. Czat zapamięta najczęściej powtarzające się informacje o nas, nasze zainteresowania, preferencje oraz kontekst projektów. W części czatów możemy przeglądać zapisane informacje w ustawieniach lub zapytać w nowym czacie: „Co o mnie pamiętasz?” i polecić usunięcie zbędnych faktów.

[12:30] Na slajdzie przedstawiono przykładowe podsumowanie pamięci profilowej autorki z ustawień ChatGPT, podzielone na sekcje: Przegląd, Praca i dydaktyka, Technologie i projekty oraz Hobby i styl spędzania czasu (zawierające m.in. informacje o pracy na UKW, tematach badań BCI/EEG/VR oraz o kocie Narkanie rasy Neva Masquerade).

Gdzie te informacje mają znaczenie? Przygotowując kolejny webinar czy wydarzenie popularyzujące naukę, nie muszę od nowa opisywać czym się zajmuję i gdzie pracuję, bo to wszystko jest zapamiętane. W życiu prywatnym nie muszę też tłumaczyć za każdym razem, że mój kot jest trenowany i spędza ze mną czas na zewnątrz.

[00:36:34] Halucynacje AI i weryfikacja odpowiedzi

Czym są halucynacje modelu AI? Halucynacja to zmyślona informacja przekazana z pełnym przekonaniem. AI jest w stanie wymyślić cały artykuł, nieistniejący paragraf czy nazwisko. Dzieje się tak, ponieważ model dąży do płynności wypowiedzi i generuje słowa na podstawie prawdopodobieństwa, a nie bazy faktów.

Punkty wzmożonej czujności (red flags): zbyt duże uogólnienia, odpowiedzi wymijające, podawanie numerów artykułów bez ich cytowania oraz brak konkretnej podstawy prawnej (sformułowania typu „zgodnie z przepisami”). Należy wtedy poprosić model o podanie konkretnego źródła i samodzielnie to zweryfikować.

Optymalizacji czasu służy zasada 80/20. Około 80% pracy wykonawczej (szkic, formatowanie, usunięcie błędów) wykonuje AI, a pozostałe 20% to praca eksperta – nasza weryfikacja i akceptacja treści. Nigdy nie używamy treści wygenerowanych bez ich przeczytania. Warto dopytywać: *„Na podstawie którego punktu dokumentu tak twierdzisz?”* – model często wtedy koryguje odpowiedź i przeprasza za błąd.

[00:40:24] Kiedy warto korzystać z AI, a kiedy nie?

AI ma odciążać, a nie stawać się kolejnym obowiązkiem. Jeśli czynności są dla nas bezproblemowe i robimy je bardzo szybko, nie włączamy AI. Przy krótkich mailach wewnętrznych ze współpracownikami warto zachować ludzką formę. Natomiast przy skomplikowanych pismach, odmowach wymagających asertywności i formalnego poukładania myśli, warto wykorzystać AI. Nawet jeśli samo promptowanie zajmie na początku dużo czasu, jest to opłacalna inwestycja w naukę formułowania wymagań i edukację czatu. Podobnie jak w Excelu – raz poświęcony czas na zbudowanie tabeli i formuł procentuje przez cały rok.

[00:45:59] Zadanie praktyczne dla uczestników

Mam dla Państwa zadanie na przełamanie pierwszej bariery: pomyślcie o jednym powtarzalnym zadaniu rutynowym wykonywanym w każdym tygodniu. Spróbujcie wykonać je wspólnie ze sztuczną inteligencją. Jeżeli uda się osiągnąć zadowalający rezultat, zapiszcie prompt w notatniku i powtórzcie zadanie w kolejnym tygodniu, sprawdzając oszczędność czasu.

[00:47:04] Wybór narzędzia: ChatGPT, Google Gemini, Claude

Czat jest głównym i najważniejszym asystentem o niskim progu wejścia. Obsługa odbywa się w przeglądarce i wymaga jedynie założenia konta. Do pracy administracyjnej polecam 3 narzędzia: ChatGPT, Google Gemini oraz Claude.

[12:30] Na slajdzie przedstawiono tabelę porównawczą trzech modeli:

- **Claude:** Główna zaleta to wzorcowa polszczyzna urzędowa, elegancki i dyplomatyczny styl wypowiedzi. Idealny do pism formalnych, decyzji i wniosków.
- **ChatGPT:** Główna zaleta to analiza danych i sprawozdawczość, elastyczny i bezpośredni styl. Świetny do ankiet, quizów, tabel i wyliczeń. Dostępność w wersji

darmowej bywa ograniczona szybszym zużywaniem tokenów w funkcjach zaawansowanych.

- **Google Gemini:** Główna zaleta to integracja z chmurą i dokumentami Google, szybkość generowania grafik oraz konkretny, wypunktowany styl. Szeroki darmowy dostęp bez ostrych podziałów na funkcje Pro.

Warto przetestować każde z tych narzędzi samodzielnie i wybrać jedno, z którym pracujemy najbardziej komfortowo.

[00:55:47] Limity i subskrypcje

Na start nie warto kupować subskrypcji – przez pierwsze miesiące warto skupić się na nauce w wersji darmowej. Zakup płatnego pakietu powinien wynikać z tego, że narzędzie realnie usprawnia naszą pracę, a darmowe limity przestają wystarczać. Monit o zużyciu tokenów pojawia się najczęściej przy zadawaniu wielu pytań w jednym oknie konwersacji w krótkim czasie, analizowaniu bardzo długich treści, analizie grafik oraz generowaniu obrazów.

[00:57:40] Praca z grafiką i generowanie plakatu

Generowanie grafik w czacie stanowi przydatny dodatek wizualny (np. ilustracja do ogłoszenia o konferencji czy synteza informacji do maila). Przykład zastosowanego przeze mnie prompta:

„Przygotuj grafikę na stronę uczelni oraz do maila dla pracowników informującą o seminarium naukowym, które odbędzie się 10 października o 11:00 w auli Copernicanum UKW. Seminarium pod tytułem: »Czy AI zawładnie światem?« wygłosi mgr inż. Jan Kowalski. Wykorzystaj logo oraz kolorystykę zgodną ze stroną UKW.”

[12:30] Na slajdzie przedstawiono wygenerowany plakat ze szkicem budynku Copernicanum w tle i poprawnie zamieszczonymi danymi tekstowymi. Autorka wskazuje na błąd w logotypie – model wstawił stary herb zamiast nowego logo UKW. Wskazuje, że dla bezpieczeństwa lepiej dołączać plik z właściwym logo jako załącznik.

[01:00:31] Iteracyjna poprawa tekstów i zmiana tonu

Niemal niemożliwym jest napisanie od razu idealnego prompta. Celem jest przerzucenie chaotycznych myśli do czatu, a rolą AI jest przekształcenie ich w elegancką polszczyznę. W promptowaniu nie trzeba martwić się o błędy ortograficzne czy interpunkcyjne.

Najczęstsze polecenia modyfikujące wygenerowany tekst:

- Zmiana tonu: *„Zmień to pismo na bardziej stanowcze”, „Przeredaguj wiadomość, aby była bardziej dyplomatyczna / łagodna”.*
- Manipulowanie długością: *„Zredukuj tekst o połowę, zachowując najważniejsze fakty”* lub *„Podaj więcej szczegółów / dopytaj mnie o brakujące dane”.*

- Korekta edytorska: wyłapywanie błędów i narzucanie zakazów edytorskich (np. *„Nie używaj długich myślników, wielkich liter w środku zdania ani dwukropków”*).

[01:05:04] Organizacja czatów i ciągłość kontekstu

Bardzo istotne jest prowadzenie tej samej sprawy w tym samym wątku konwersacji. Cząty można przypinać u góry oraz zmieniać ich nazwy. Dzięki temu po tygodniu czy miesiącach wracamy do tego samego czatu (np. sprawozdań z Rady Kierunku), zachowując ciągłość kontekstu. Nową rozmowę zaczynamy tylko wtedy, gdy przechodzimy do całkowicie nowego zadania lub gdy stary czat uległ przeładowaniu.

MODUŁ II: Przetwarzanie wiedzy i tworzenie dokumentacji – praca z własnymi plikami i interaktywnym edytorem AI

[01:08:31] Wprowadzenie do drugiego modułu

W ramach drugiego modułu poruszymy zagadnienia pracy z własnymi plikami, syntezy informacji, budowania bezpiecznej bazy wiedzy, obsługi interaktywnego edytora AI (Canvas), selekcjonowanej edycji, pracy na tabelach oraz ostatecznej korekty i eksportu dokumentacji.

[01:09:56] Praca z załącznikami (zasilanie wiedzą)

Zamiast budować bardzo długi prompt, warto dorzucić pliki źródłowe: regulaminy, notatki wielostronicowe, PDF-y, a nawet skany czy zdjęcia pism. AI bez problemu dokona syntezy z wielu źródeł naraz. Wgrywanie własnych plików zapewnia aktualność wiedzy i 100% kontroli nad źródłem, eliminując ryzyko pobrania przez model nieaktualnych rozporządzeń z Internetu.

[01:12:51] Przesłuchiwanie dokumentów i zasada Grounding (Zero Halucynacji)

Zamiast samodzielnie przeszukiwać kilkadziesiąt stron dokumentu, wgrywamy go do czatu i wymuszamy odpowiedź wyłącznie na bazie treści z pliku. Stosujemy wtedy metodę ścisłego zakotwiczenia (Grounding).

Przykładowe klauzule do prompta:

- *„Wczytaj załączony plik Regulamin pracy UKW. Odpowiedz na pytanie, jakie są warunki i limity rozliczania godzin ponadwymiarowych dla asystentów dydaktycznych. Ważne: korzystaj wyłącznie z danych zawartych w załączonym pliku. Jeśli w dokumencie nie ma tej informacji, napisz wprost: brak w źródle.”*
- *„Do każdej podanej informacji podaj numer strony lub paragraf z pliku. Nie uzupełniaj odpowiedzi swoją wiedzą ogólną z Internetu.”*

[01:14:56] Trzy poziomowe streszczenie z jednego źródła

Jednym promptem w jednym momencie możemy przygotować 3 różne komunikaty na bazie jednego źródła wiedzy:

1. **Poziom 1 (Mail do pracowników):** Najbardziej szczegółowy i rozbudowany komunikat z wytycznymi krok po kroku i terminami.
2. **Poziom 2 (Komunikat na stronę www):** Zwięzły, oficjalny skrót najważniejszych zmian.
3. **Poziom 3 (Notatka na social media):** Maksymalny skrót w pigułce, lekki język, punktatory, kluczowe daty i emoji.

[01:16:39] Trudne pliki i porównywanie dokumentów

Przy skanach i zdjęciach dokumentów AI wykorzystuje technologię OCR. W pierwszym etapie warto poprosić o weryfikację czytelności tekstu (np. zapytać co zawarte jest na skanie i sprawdzić poprawność). Przy wielostronicowych tabelach w PDFie należy uważać, czy model nie pogubił nazw kolumn. W miarę możliwości lepiej załączać pliki w formacie DOCX lub XLSX.

AI doskonale sprawdza się przy porównywaniu dokumentów (np. dwóch wersji uchwały czy ofert przetargowych). Przykład prompta: „*Wczytaj plik pierwszy (stary) i plik drugi (nowy). Przygotuj tabelę różnic ze wskazaniem co zmieniono, co usunięto i co dodano nowego.*”

[01:20:18] NotebookLM – darmowa baza wiedzy od Google

NotebookLM (notebooklm.google.com) to darmowe narzędzie od Google dedykowane do pracy z załącznikami (do 50 dokumentów w jednym notatniku). Różni się od standardowego czatu tym, że działa w środowisku zamkniętym (grounding) – odpowiada wyłącznie na podstawie wgranych materiałów, a w przypadku braku danych wprost informuje o braku w źródle.

Zastosowania w administracji i dydaktyce: budowanie baz regulaminów, obsługa wydarzeń (programy, ustalenia), porządkowanie własnych materiałów dydaktycznych. W interfejsie narzędzia po wgraniu źródeł można jednym kliknięciem generować mapy myśli, infografiki, tabele, prezentacje, fiszki oraz podsumowania w formacie audio lub video.

[12:30] Prezentacja wygenerowanego automatycznie przez NotebookLM podsumowania video oraz audycji audio (podcastu) na podstawie wgranego Regulaminu Studiów UKW oraz instrukcji logowania do USOS Web.

[01:31:47] Higiena źródeł i cytowania w NotebookLM

Przy wgrywaniu plików do NotebookLM należy dbać o precyzyjne nazewnictwo z datą i wersją oraz na bieżąco usuwać nieaktualne dokumenty. Każde stwierdzenie wygenerowane przez NotebookLM posiada numerowany odnośnik. Kliknięcie przypisu przenosi nas bezpośrednio do dokładnego fragmentu i strony w naszym pliku źródłowym, co umożliwia natychmiastową weryfikację cytatu prawnego. Narzędzie idealnie nadaje się do tworzenia wewnętrznych baz FAQ dla sekretariatów i oszczędza czas przy wdrażaniu nowych pracowników.

[01:38:34] Porównanie: Standardowy czat vs NotebookLM

- **Standardowy czat (ChatGPT, Gemini, Claude):** Przeznaczony do szybkiego pisania, redagowania, burzy mózgów i krótkich pytań. Bazuje na wiedzy ogólnej modelu plus pojedynczych załącznikach.
- **NotebookLM:** Przeznaczony do analizy wielu dokumentów (do 50 plików) i budowania stałych baz przepisów. Bazuje wyłącznie na wgranych źródłach (100% grounding), zapewniając klikalne odnośniki do fragmentów dokumentów. (Standardowy czat to motyl na szybkie pytania, a NotebookLM to armata na skomplikowane analizy prawne).

[01:41:22] Interaktywny edytor AI (Gemini Canvas)

Gemini Canvas to interaktywna przestrzeń robocza z podzielonym ekranem: po lewej stronie znajduje się okno dyskusji z czatem, a po prawej żywy dokument (tekst, kod, prezentacja, tabela, quiz). Podobne rozwiązania to Canvas w ChatGPT oraz Artifacts w Claude.

Uruchamiamy go przyciskiem plusa przy oknie prompta. Przykładowy prompt: *„Utwórz dokument, w którym przygotujesz szkielet sprawozdania z rady kierunku Socjoinformatyka. Dokument musi zawierać datę spotkania, skład rady, listę obecności, agendę, przebieg oraz listę załączników.”*

[01:45:05] Zalety pracy w środowisku Canvas

- Dokument znajduje się stale w prawym panelu (brak konieczności przeszukiwania długiego czatu).
- Edycja selekcyonowana: zaznaczamy myszką konkretny akapit i prosimy o zmianę tylko tego fragmentu – reszta tekstu pozostaje nienaruszona.
- Wersjonowanie: możliwość łatwego cofania zmian do poprzednich wersji.
- Zachowanie formatowania i eksport: bezpośredni zapis do Google Docs, Worda (DOCX), Excela (XLSX) lub PDF.

[01:48:32] Zasady tworzenia dokumentów w Canvas

1. **Od ogółu do szczegółu:** Tworzymy najpierw szkielet i spis treści, a dopiero potem uzupełniamy poszczególne sekcje.
2. **Stosowanie znaczników (placeholders):** Umieszczamy w klamrach oznaczenia takie jak [Data posiedzenia] czy [Nazwisko studenta], aby uzupełnić dane poufne lokalnie poza edytorem AI.
3. **Pętla iteracyjna (tryb hybrydowy):** AI tworzy szkic -> człowiek nanosi ręczne poprawki w dokumencie -> zaznaczamy fragment i prosimy AI o wyszlifowanie stylu lub korektę błędów.
4. **Zapisywanie kopii roboczych:** Po osiągnięciu ważnego etapu warto pobrać plik lokalnie na komputer jako zabezpieczenie przed przypadkowym nadpisaniem.

[01:55:32] Edycja selekcyonowana i przyciski szybkich akcji

Zaznaczenie kursorem myszy fragmentu tekstu w prawym panelu otwiera okienko prompta

inline (z ikoną ołówka i gwiazdki), pozwalając wydać komendę odnoszącą się tylko do zaznaczenia (np. „Przekształć akapit na język bardziej oficjalny”, „Zmień stronę bierną na aktywną”).

Dodatkowo dostępne są 3 przyciski szybkich akcji:

1. *Zmień długość*: 5-stopniowa skala skracania lub wydłużania tekstu.
2. *Zmień ton*: 5-stopniowa skala dostosowania stylu (od swobodnego do formalnego).
3. *Zasugeruj zmiany*: Błyskawiczna poprawa literówek, interpunkcji i spójności gramatycznej.

[01:59:34] Praca na tabelach i modułowe sprawozdania

Canvas zapobiega rozjeżdżaniu się kolumn i umożliwia łatwe dodawanie nowych kolumn poleceniem w czacie. Pozwala przekształcać nieustrukturyzowany tekst (maile, protokoły) w czyste tabele oraz agregować dane do sprawozdań.

Metoda modułowa tworzenia sprawozdań rocznych obejmuje 4 kroki: 1.

Generowanie spisu treści; 2. Wypełnienie Rozdziału 1 na podstawie Pliku A; 3.

Wypełnienie Rozdziału 2 na podstawie Pliku B; 4. Synteza, wnioski i ostateczny prompt korygujący styl i formy grzecznościowe.

[02:03:35] Checklista weryfikacyjna przed eksportem (Human-in-the-Loop)

Zgodnie z zasadą Human-in-the-Loop za dokument podpisany własnym nazwiskiem zawsze odpowiada człowiek. Przed eksportem należy zweryfikować:

1. Fakty i dane (czy liczby, daty i nazwiska są prawdziwe).
2. Brak halucynacji (czy przytoczone paragrafy i źródła istnieją).
3. Kontekst i sens (czy treść jest spójna i logiczna).
4. Kompletność (czy nie brakuje sekcji lub załączników).
5. Poprawność językową.
6. Formatowanie i układ graficzny.
7. Dostępność cyfrową według standardu WCAG (jeżeli dokument będzie publikowany cyfrowo).

MODUŁ III: Własny asystent AI z automatyzacją procesów, personalizacja narzędzi i praca z nagraniami

[02:06:06] Wprowadzenie i porównanie: Ogólny czat vs Dedykowany asystent

W trzecim module omówimy tworzenie własnych asystentów AI (Gems w Gemini, GPTs w ChatGPT, Projects w Claude), personalizację instrukcji, pracę z mową i nagraniami audio, łączenie narzędzi w łańcuchy pracy (workflow) oraz podsumowanie dobrych nawyków.

Ogólny czat posiada wiedzę ogólną i pamięć profilową, ale przy mieszaniu wątków trudno utrzymać w nim sztywny szablon edytorski dla konkretnego zadania.

Dedykowany asystent (Gem / GPT / Project) to wyizolowane środowisko stworzone pod jeden konkretny proces (np. tworzenie sprawozdań z Rady Kierunku). Posiada

na stałe zdefiniowaną rolę, instrukcje, zakazy oraz pliki wzorcowe. Konfiguracja takiego bota zajmuje około 5 minut.

[02:11:12] Kiedy warto stworzyć asystenta i współdzielenie w zespole
Asystenta warto utworzyć przy cyklicznych procedurach (sprawozdania, protokoły, opisy PKA, odpowiedzi na wnioski), eliminując potrzebę ciągłego wpisywania długich wytycznych. Gotowego asystenta można udostępnić zespołowi lub sekretariatowi za pomocą linku, co zapewnia jednolity standard korespondencji i oszczędność czasu dla całego biura.

[02:13:13] Proces konfiguracji Gema w Google Gemini
W panelu bocznym wybieramy opcję „Gemy” -> „Nowy gem”. Wypełniamy pola: Nazwa, Opis, Instrukcje (główny prompt), Załączniki (wzorcowe pliki i regulaminy) oraz opcjonalnie Narzędzie domyślne (np. edytor Canvas lub generowanie grafik).

[02:14:38] Budowanie zaawansowanej instrukcji (Few-shot prompting i sekcje)

- Precyzyjne określenie roli i celu (np. *„Jesteś redaktorem sprawozdań z posiedzeń rady kierunku”*).
- Stosowanie sekcji z nagłówkami z hasztagami (#Rola, #Zasady, #Wyjście) w oknie instrukcji.
- Wykorzystanie Few-shot prompting: wklejenie w instrukcji kilku przykładowych par wejście-wyjście (wzorcowych pism), co wymusza idealny format odpowiedzi.
- Określenie zakazów i nakazów: np. *„Nigdy nie stosuj ozdobnych wstępów (np. 'Oto wygenerowane sprawozdanie')”, „Unikaj myślników i zbędnych wypunktowań”, „Unikaj słów marketingowych (kluczowe, pigułka, dynamiczne)”*.

[02:21:44] Cztery przykłady zastosowania dedykowanych asystentów

1. **Asystent do korespondencji e-mail:** Przekształca surowe wypunktowania w eleganckie maile z dbałością o zwroty grzecznościowe.
2. **Asystent do ścisłej korekty:** Skonfigurowany wyłącznie na usuwanie błędów i interpunkcji, z kategorycznym zakazem modyfikacji treści merytorycznej.
3. **Prywatny asystent balkonowy (przykład autorki):** Posiada w bazy wiedzy spis roślin, pojemników i nawozów; pozwala zadawać szybkie pytania głosowe ze zdjęciem liścia bez opisywania warunków od nowa.
4. **Asystent dydaktyczny ds. sylabusów:** Posiada wgrane wytyczne PKA i matryce efektów uczenia się, weryfikując spójność punktów ECTS i opisów przedmiotów.

[02:25:10] Praca z mową i tryb głosowy (Voice Mode)

Mówimy średnio 3 razy szybciej niż piszemy na klawiaturze (130–150 słów/minutę vs 40 słów/minutę). Nowoczesne modele świetnie radzą sobie z językiem polskim i odmianą. W

trybie Voice możemy podyktować chaotyczne myśli, a model wygeneruje z tego uporządkowaną notatkę lub oficjalnego maila (np. „*Napisz maila do Anny... a nie, czekaj, do dr Nowak w sprawie protokołu, powiedz, że dostaniemy go w piątek, ale niech przyspieszy, bo pani dziekan czeka*”).

[02:27:39] Pełny zapis rozmowy na żywo w trybie głosowym z ChatGPT

(Poniżej przedstawiono pełną, dosłowną transkrypcję konwersacji głosowej przeprowadzonej w czasie rzeczywistym podczas szkolenia):

Adrianna Piszcz: Czat GPT, jesteśmy teraz na żywo na szkoleniu AI dla administracji. Streść w dwóch zdaniach, o czym się tutaj właściwie teraz uczymy.

ChatGPT: Jasne! Uczycie się dziś, jak praktycznie wykorzystywać AI w codziennych zadaniach administracyjnych – przede wszystkim w pisaniu i porządkowaniu tekstów, tworzeniu i ulepszaniu promptów oraz rozumieniu ograniczeń tych narzędzi. I co ważne, uczymy się też podejścia odpowiedzialnego, czyli weryfikowania wyników i dbania o dane, żeby AI wspierało pracę zamiast ją komplikować.

Adrianna Piszcz: Tak, zgadza się. A gdybyś miał powiedzieć, jaka jest jedna, najważniejsza rzecz, którą osoba pracująca w administracji powinna wynieść z tego szkolenia?

ChatGPT: Żeby traktować AI jako mocnego asystenta, a nie nieomylny autorytet. Czyli używać, ale zawsze weryfikować, wiedzieć, kiedy nie ufać ślepo, dbać o dane i mieć nawyk zadawania precyzyjnych pytań. To naprawdę zmienia efekty.

Adrianna Piszcz: Oki, to podaj mi 3 konkretne przykłady zadań w administracji, przy których AI może realnie oszczędzać czas – być pomocnikiem, a nie wrogiem.

ChatGPT: Po pierwsze: redagowanie i poprawianie pism oraz e-maili. AI przygotowuje szkic odpowiedzi, skróci zbyt długą wiadomość albo zrobi bardziej formalny ton. Po drugie: streszczanie dokumentów. Długi regulamin czy raport – AI wyłuska kluczowe punkty i szybciej zorientujesz się w treści. Po trzecie: tworzenie prostych materiałów roboczych – planu spotkania, checklisty, tabelki z zadaniami. I ważne: AI tu nie zabiera pracy, tylko odciąża od rutyny, a decyzje i odpowiedzialność dalej są po stronie człowieka.

Adrianna Piszcz: Ok, a które z tych przykładów uważasz, że da największą oszczędność czasu i dlaczego tak uważasz?

ChatGPT: Największą oszczędność zwykle daje redagowanie i poprawianie tekstów, bo to pojawia się niemal codziennie w wielu działach. AI może robić pierwszy szkic albo edytować, a człowiek sprawdza fakty i finalną wersję. To się najbardziej skaluje.

Adrianna Piszcz: Super, to może prześledźmy sobie taki konkretny przykład: dostałam długiego i chaotycznego maila od pracownika. Jak mógłbyś mi pomóc przygotować odpowiedź tak, żeby była rzeczowa, uprzejma, krótka, ale jednocześnie ludzka?

ChatGPT: Zamiast od razu sama pisać, możesz wkleić treść, a ja wyłuskam kluczowe kwestie i przygotuję krótki, uprzejmy szkic odpowiedzi. Jeśli będzie zbyt formalny albo za bardzo skrótowy, dopasujemy ton, żeby pasował do stylu komunikacji w Twojej instytucji. Ty decydujesz, co ostatecznie wychodzi na zewnątrz.

Adrianna Piszcz: Hej, a co powinien zrobić człowiek zanim kliknie „wyślij”?

ChatGPT: Powinien przeczytać wszystko własnymi oczami i sprawdzić 3 rzeczy: czy fakty i liczby się zgadzają, czy ton jest odpowiedni i czy nie ma danych wrażliwych. I dopiero potem wysłać. AI pomaga z brudnopisem, ale odpowiedzialność za finalną wersję zawsze zostaje po stronie człowieka.

Adrianna Piszcz: Zgadza się, masz rację, to bardzo ważne – Human-in-the-Loop. A powiedz mi, czy musimy umieć w jakiś specyficzny sposób pisać prompty, żeby dobrze pracować z AI?

ChatGPT: Nie musicie mieć żadnych czarów z promptami. Często wystarczy proste polecenie, a potem doprecyzowanie krok po kroku: „Skróć to do 3 zdań”, „Zrób bardziej przyjacielski ton”. Im jasno mówisz czego chcesz, tym lepiej. A jeśli podasz kontekst – na przykład, że to do komunikacji wewnętrznej – odpowiedź zwykle będzie trafniejsza. To bardziej rozmowa niż zaklęcia, więc spokojnie pytaj i doprecyzowuj.

Adrianna Piszcz: Ok, to pokaż na przykładzie różnicę między takim ogólnym poleceniem, które sobie naturalnie wpisujemy, a takim dobrze przygotowanym, stronicowym promptem.

ChatGPT: Pewnie! I tutaj wiele osób się zaskakuje, bo różnica nie musi być aż tak duża. Proste: „Napisz odpowiedź na tego maila” już działa. Ale dopracowane polecenie doda szczegóły: trzymaj ton uprzejmy i profesjonalny, około 150 słów, odpowiedz na wszystkie pytania, a jeśli czegoś nie wiesz, zaznacz to zamiast zmyślać. I to zwykle już daje świetny efekt. Ale nie musisz od razu pisać stronicowych punktów. W praktyce zacznij prosto, potem doprecyzuj. To właśnie największy plus pracy z AI: możesz prowadzić dialog zamiast próbować od razu ułożyć idealne polecenie na start.

Adrianna Piszcz: No bo właściwie to nie da się chyba tak od razu być idealnym, prawda, w tym promptowaniu?

ChatGPT: Dokładnie! Nikt nie zaczyna idealnie. Traktuj to jak naukę rozmowy: jedno proste polecenie, potem poprawka, potem jeszcze jedna – z czasem robi się to naturalne.

Adrianna Piszcz: No dobrze, no to, skoro jest ogólnie pomocna, to czego zdecydowanie nie powinniśmy powierzać bezrefleksyjnie AI?

ChatGPT: Nie powierzaj bezrefleksyjnie danych osobowych, informacji poufnych, treści wrażliwych prawnie ani niczego, co wymaga zgodności z przepisami bez sprawdzenia zasad instytucji. I nie zostawiaj AI samej z interpretacją. W takich sprawach liczy się człowiek, przepisy i polityka organizacji dużo bardziej niż sama wygoda.

Adrianna Piszcz: Czyli jeśli wygenerujesz mi odpowiedź, która brzmi bardzo profesjonalnie, to nie musi znaczyć, że ona jest prawdziwa?

ChatGPT: Dokładnie! Profesjonalne brzmienie nie gwarantuje prawdy. Model potrafi mówić bardzo przekonująco, a jednak się mylić. Dlatego weryfikacja jest obowiązkowa, zwłaszcza przy ważnych sprawach. Najlepiej traktować AI jako pierwszy szkic, nie wyrocznię.

Adrianna Piszcz: No dobra, to wtedy co z danymi, które mam w pracy? Czy mogę wszystko wklejać do czatu, nawet rzeczy, które mają dane osobowe czy informacje poufne?

ChatGPT: To zależy od narzędzia i polityki organizacji. Najbezpieczniej nie wklejać danych osobowych ani poufnych do publicznych narzędzi AI, chyba że macie wyraźne zasady i zgodę na takie użycie. Jeśli trzeba poprosić AI o pomoc, rozważcie anonimizację danych albo użycie narzędzi firmowych z odpowiednimi zabezpieczeniami. Najlepiej najpierw sprawdzić zasady waszej instytucji, a dopiero potem decydować co wolno wkleić. To pozwoli wam korzystać z tych narzędzi mądrze i bezpiecznie.

Adrianna Piszcz: Ja myślę, że ta anonimizacja to taki must-have w pracy, nie należy tego traktować jako coś opcjonalnego.

ChatGPT: Zgadzam się. Anonimizacja to dobra praktyka, nawet jeśli polityka instytucji dopuszcza użycie konkretnych narzędzi. Im mniej wrażliwych danych wklejacie i udostępniacie, tym mniejsze ryzyko i więcej spokoju.

Adrianna Piszcz: No dobrze, no to ostatni temat. Jak prześledzisz całą agendę naszego szkolenia, to jaki temat twoim zdaniem może być najtrudniejszy dla początkującego użytkownika AI?

ChatGPT: Myślę, że dla większości początkujących najtrudniejsze nie jest samo AI, tylko promptowanie, czyli nauczanie się rozmawiać. Wszyscy myślą, że muszą od razu pisać idealne polecenia, a potem się frustrują.

Adrianna Piszcz: No dobrze, to jakbyś miał wytłumaczyć takiej osobie, która dopiero zaczyna, w jaki sposób ma w takim razie tworzyć te swoje pierwsze prompty?

ChatGPT: Powiedziałbym: nie traktuj promptów jak programowania, myśl o nich jak o rozmowie z kompetentnym współpracownikiem. Wystarczą 4 proste elementy: co chcesz osiągnąć, jaki jest kontekst, jaki efekt i czego oczekujesz od odpowiedzi. Nie próbuj od razu napisać perfekcyjnego prompta. Jeśli pierwsza odpowiedź nie pasuje – doprecyzuj, skróć, dodaj konkretów, napisz prościej albo zmień ton na bardziej formalny. Tak wygląda codzienna praca z AI krok po kroku przez dopasowywanie odpowiedzi.

Adrianna Piszcz: Dobrze, to na zakończenie chciałabym, żebyś dał złotą radę dotyczącą korzystania z AI w pracy w administracji dla naszych uczestników szkolenia.

ChatGPT: Niech AI będzie waszym partnerem i skrótem. Pytajcie, doprecyzowujcie, sprawdzajcie i nie bójcie się powiedzieć: „nie, to nie to”. Najlepsze efekty przychodzą w rozmowie, nie w pierwszym podejściu.

Adrianna Piszcz: Super, dziękujemy bardzo za rozmowę.

ChatGPT: Cała przyjemność po mojej stronie, powodzenia na szkoleniu!

[02:39:01] Praca z plikami audio i RODO

Gotowe nagrania z dyktafonu w telefonie (MP3, MP4) można wgrywać jako źródło do czatu, asystenta lub NotebookLM i wykonać automatyczną transkrypcję mowy na tekst lub wyciągnąć tabelę zadań, terminów i osób odpowiedzialnych.

Obowiązują tu bardzo ważne zasady formalne i RODO: zawsze informujemy uczestników o nagrywaniu i uzyskujemy ich zgodę. Obowiązuje zakaz wgrywania nagrań zawierających wrażliwe dane osobowe (PESEL-e, dane medyczne) – takie fragmenty

należy wyciąć z pliku audio przed wgraniem do chmury. Przy trudnym lub zaszumionym dźwięku stosujemy prompt korygujący: „*Rozpoznaj głównych rozmówców jako Osoba A i Osoba B. Ignoruj szumy i przejęzyczenia. Jeśli pojęcie brzmi jak skrót akademicki, dopasuj go do kontekstu szkolnictwa wyższego.*”

[02:43:17] Łańcuchy pracy (AI Workflow) i przykłady map procesów

Bywają przypadki, gdy nie używamy jednej aplikacji do wykonania całego zadania. Musimy umieć podzielić pracę na mniejsze zadania i przekazywać je do dedykowanych narzędzi (pomocna wyszukiwarka narzędzi: *There's An AI For That* – aiforthat.com [DO WERYFIKACJI: „theresanaiforthat.com”]).

Przykłady łańcuchów procesów:

- *Od nagrania do pisma*: dyktowanie notatki na telefonie -> przekazanie tekstu do Asystenta Stylu -> przeniesienie do edytora Canvas na ostateczny eksport.
- *Od Regulaminu do decyzji*: weryfikacja skomplikowanego wniosku studenta w NotebookLM z odniesieniem do paragrafu -> przekazanie ustaleń do Asystenta Mailowego w celu wygenerowania odpowiedzi.
- *Poranna korespondencja*: dyktowanie myśli na głos przy kawie -> Asystent Mailowy -> gotowa wiadomość w kilka sekund.

[02:47:07] Zmiana roli pracownika i podsumowanie dobrych nawyków

Standardowy typ pracy to pisanie od zera i ręczne przeglądanie regulaminów.

Zmieniony tryb pracy z AI czyni z pracownika zarządcę i redaktora (editor-in-chief): narzędzia tworzą pierwszy szkic, a człowiek weryfikuje, poprawia i podejmuje decyzje (zasada Human-in-the-Loop).

Złote zasady higieny pracy z AI:

1. **Anonimizacja**: bezwzględne usuwanie PESELI, imion i nazwisk.
2. **Weryfikacja**: zawsze sprawdzamy daty, liczby i przytoczone paragrafy.
3. **Dobór narzędzia**: dopasowanie narzędzia do skali problemu (Czat / NotebookLM / Canvas / Gemy).
4. **Strategia małych kroków**: stopniowa automatyzacja (Tydzień 1: Zwykły czat do korekty; Tydzień 2: NotebookLM; Tydzień 3: Dedykowany Gem; Tydzień 4: Praca z głosem).

[02:53:10] Realna oszczędność czasu i zakończenie

Realny zysk czasowy: analiza 50-stronicowego PDF-a ulega skróceniu z 2 godzin do kilku minut; pisanie pisma lub maila z 20 minut do 2 minut; synteza nagrania ze spotkania z kilkudziesięciu minut do 5 minut. Praca z AI przypomina jazdę na rowerze – wymaga praktyki, ale z każdym dniem staje się bardziej intuicyjna.

[02:54:31] Powtórzenie informacji projektowych i licencyjnych. Prowadząca Adrianna Piszcz dziękuje za uwagę i życzy owocnych testów i pracy ze sztuczną inteligencją.

[02:55:35] Pełny tekst piosenki końcowej (wygenerowanej w Suno AI)

(Na zakończenie szkolenia odtworzono utwór muzyczno-słowny wygenerowany przez autorkę w narzędziu Suno AI):

*Zalogowani przed ekranem, każdy w swoim biurze siedzi
Wokół maile, sprawozdania, sterta spraw i trudnych odpowiedzi
Wchodzimy w ten webinar z głową pełną pytań i obaw
Czy ten cały świat AI to dla nas odpowiednia droga?
Ale piszemy pierwsze prompty, w oknie czatu rośnie szkic
I nagle wiemy, że z asystentem nie trzeba się bać nic!*

*Krok po kroku, z każdym slajdem Canvas, Gemy, audio w tle
Odnajdziemy w tym spokój
Praca w biurze lżejsza staje się!*

[Refren]

*Uczyliśmy się razem (razem, razem)
To był piękny, mądry czas!
Z AI krok po kroku
Zmienia się nasz biurowy świat!
Małe kroki w stronę zmian
Zyskasz czas i prostszy plan!*

*Ktoś pokazał inne spojrzenie, ktoś dopisał spójny wzór
Szybka synteza, czyste dane – zniknął z biurka stary mur
Było w tym tyle uważności i uczelnianej ludzkiej troski
Jakby powszedni dzień w administracji stał się prostszy*

*Krok po kroku, z każdym slajdem Canvas, Gemy, audio w tle
Odnajdziemy w tym spokój
Praca w biurze lżejsza staje się!*

[Refren] ...

*To nie AI zmienia świat, lecz człowiek z nową wiedzą
Gdy odwagi ma trochę więcej i pomysły w nim siedzą
Dziś wiemy więcej, zyskujemy czas na to, co ważne jest
Otwarta głowa i mądre prompty – to nasz wspólny gest!*

[Refren] ...

*Dziękuję Wam za ten czas
Za małe kroki i odwagę w nas
Powodzenia w pracy z AI
Niech ten dobry powiew trwa!*