



Dokumentacja Użytkownika Systemu

AsWiseAI

Platforma analizy i zarządzania wiedzą w oparciu o sztuczną inteligencję
Konwersacyjna, Inteligentna Wyszukiwarka AI

Producent:

Firma Informatyczna EnterSoft

AsWiseAI.pl, EnterSoft.pl

Wersja v 2.05 <-> 2026

1.	Wprowadzenie do AsWiseAI.....	4
	Co to jest AsWiseAI?	4
	Bezpieczeństwo i poufność danych	5
	Wymagania przeglądarki	5
	Wersja mobilna.....	6
	Główne zalety systemu.....	6
2.	Rozpoczęcie pracy.....	9
	Logowanie do aplikacji	9
	Interfejs użytkownika - ogólny przegląd	9
	Pasek nawigacji bocznej (Sidebar)	10
	Pasek górny (TopBar)	11
	Szybki start.....	12
3.	Główna funkcjonalność - Baza wiedzy i analiza dokumentów	14
	1. Przesyłanie i przetwarzanie dokumentów	14
	Jak wgrać plik?	14
	Obsługiwane formaty plików (PDF, PNG, JPG, TXT)	15
	Status przetwarzania i powiadomienia.....	16
	2. Zarządzanie dokumentami	17
	Przeglądanie i filtrowanie dokumentów.....	17
	Podgląd treści i źródeł	18
	Usuwanie dokumentów.....	18
4.	Interakcja z AI	19
	1. Jak działa przepływ zapytania: Od pytania do odpowiedzi	19
	2. Główne tryby interfejsu	21
	Tryb "AI Fakt"	22
	Jak zadawać pytania?	22
	Przykłady i możliwości interakcji z LLM	22
	Tryb "AI Agent Analityczny"	24
	Rozpoczynanie i zarządzanie konwersacjami.....	25
	Prowadzenie dialogu	25
	Wykorzystanie kontekstu	25
	3. Kluczowe Korzyści Architektury Agentowej.....	26
	4. Porównanie AI Agent & AI Fakt	28
5.	Architektura Przepływu Zapytania.....	29
	Definicje Kluczowych Komponentów	30

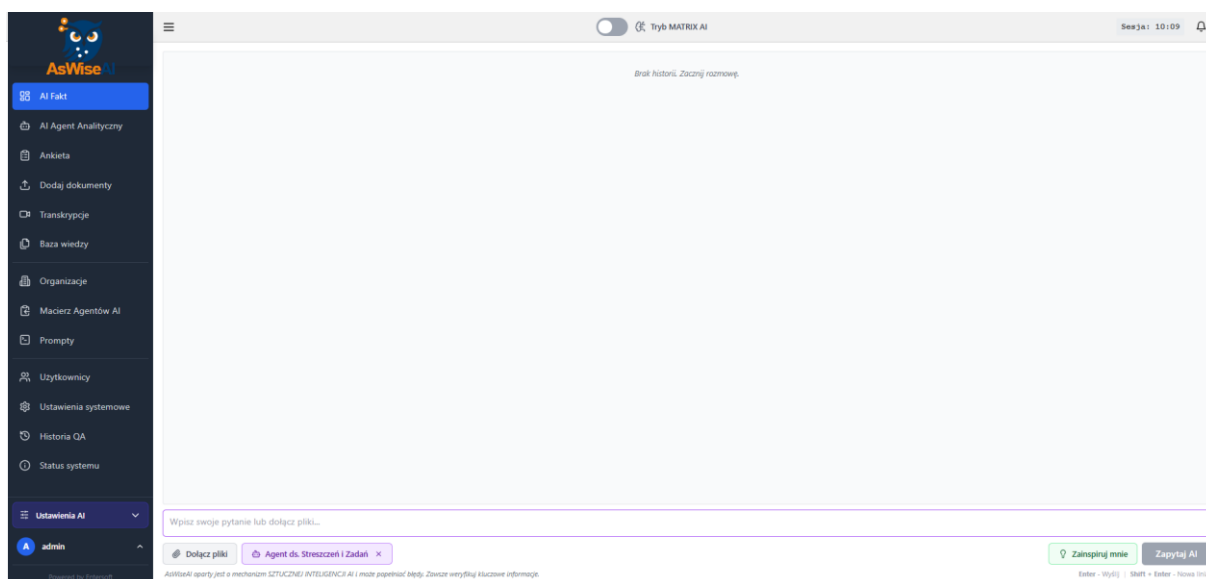
Kluczowe Zasady.....	31
Zaawansowane Możliwości	31
6. Wskazówki i inspiracje	33
1. Panel sugestii pytań AI.....	33
2. Weryfikacja informacji	33
7. Funkcje administracyjne i ustawienia	34
1. Zarządzanie kontem (Ustawienia użytkownika)	34
Zarządzanie hasłem	34
Czyszczenie historii czatów	34
Eksport historii rozmów (Eksport chatów)	34
2. Panel Administratora: Zarządzanie Główne.....	36
Zarządzanie użytkownikami i organizacjami.....	36
Macierz Agentów AI.....	36
Zarządzanie promptami.....	39
3. Panel Administratora: Ustawienia Systemowe.....	42
Automatyzacja i Przepływy Pracy (Workflow Engine)	42
Archiwizacja AI.....	55
Integracje (LDAP / Active Directory)	57
Tokeny API (Dostęp dla Aplikacji Zewnętrznych)	59
Zasady Działania i Schemat Przepływu Danych	60
Schemat Uwierzytelniania i Audytu Tokena	61
Generowanie Nowego Tokena.....	65
Unieważnianie Tokena (Revoke)	66
Zarządzanie Modelami AI	67
Zarządzanie regułami czyszczenia.....	68
Kopia zapasowa i resetowanie organizacji.....	69
Resetowanie danych.....	69
8. Panel Analityczny	70
1. Pulpit Analityczny	70
Statystyki Ocen Eksperckich	71
2. Ocena Jakości AI (Ewaluacja RAG)	71
1. Zakładka "Do Oceny"	72
2. Zakładka "Ocenione"	74
3. Aktywność Użytkowników	74
4. Baza Wiedzy.....	75

9. Status systemu.....	76
10. Pomoc i wsparcie.....	82
Najczęściej zadawane pytania (FAQ)	82
Kontakt z pomocą techniczną.....	82
Słowniczek pojęć	83
11. Rozwiązywanie problemów i dobre praktyki.....	85
Najczęstsze problemy i rozwiązania (FAQ Techniczne)	85
Dobre praktyki i optymalizacja (Pro-Tips).....	86
Praktyczny przykład: Automatyzacja faktur kosztowych	86
12. Zaawansowane Funkcje Interfejsu Czatów	88
13. Transkrypcja Plików Audio i Wideo.....	90
14. Zarządzanie Instancjami Chatbotów (Widgety WWW)	93
1. Tworzenie i Edycja Chatbota.....	93
Sekcja A: Wygląd i Tożsamość.....	94
Sekcja B: Inteligencja i Zachowanie	95
Sekcja C: Baza Sugestii	95
2. Zarządzanie na Liście i Kopiowanie Kodu	96
Jak wdrożyć bota na stronę?	96
15. Podsumowanie	99

1. Wprowadzenie do AsWiseAI



Co to jest AsWiseAI?



AsWiseAI to zaawansowana platforma cyfrowa, która pełni rolę analitycznego asystenta dokumentów, wykorzystując sztuczną inteligencję, aby pomóc Twojej firmie w przetwarzaniu i analizie treści. Została zaprojektowana z myślą o bezpieczeństwie i poufności, działając w modelu **on-premise**. Oznacza to, że całe oprogramowanie i wszystkie Twoje dane są przechowywane wyłącznie na Twojej własnej infrastrukturze, bez wysyłania ich do jakichkolwiek zewnętrznych usług w chmurze.

Zarządzanie rosnącą liczbą dokumentów i efektywne wydobywanie z nich informacji stanowi wyzwanie dla wielu organizacji. Przeszukiwanie obszernych archiwów w poszukiwaniu konkretnych danych jest często procesem czasochłonnym i mało wydajnym.

Główne przeznaczenie platformy to:

- **Wsparcie w analizie dokumentów:** AsWiseAI potrafi przetwarzać różne typy dokumentów, takie jak umowy, raporty, instrukcje czy faktury.
- **Zwiększenie efektywności:** Dzięki technologii RAG (Retrieval-Augmented Generation), system znajduje w dokumentach odpowiedzi na pytania, które zadajesz, oszczędzając Twój czas.
- **Bezpieczne i kontrolowane środowisko:** Dzięki wdrożeniu on-premise, masz pełną kontrolę nad swoimi danymi i spełniasz wymagania dotyczące poufności.

Bezpieczeństwo i poufność danych

Poufność i bezpieczeństwo Twoich danych są dla nas priorytetem. Architektura AsWiseAI została zaprojektowana, aby chronić Twoje informacje na kilka sposobów:

- **Model wdrożenia on-premise:** Wszystkie dokumenty, wektory w bazie danych oraz historia zapytań są przechowywane lokalnie, na Twojej infrastrukturze. Nie są wysyłane do żadnych zewnętrznych serwerów ani usług chmurowych, co minimalizuje ryzyko wycieku danych.
- **Hashowanie haseł:** Twoje hasła są przechowywane w formie haszowanej, co uniemożliwia ich odczytanie i zwiększa bezpieczeństwo.
- **Kontrola dostępu oparta na rolach (RBAC):** Dostęp do poszczególnych funkcji aplikacji jest zarządzany za pomocą ról i uprawnień. Zapewnia to, że tylko upoważnieni użytkownicy mają dostęp do wrażliwych danych i funkcji.
- **Izolacja danych organizacji:** Każda organizacja ma przypisaną własną, odizolowaną kolekcję w bazie danych Qdrant, co zapobiega mieszanii się danych między firmami.
- **Brak inwazyjnych plików cookie:** Aplikacja nie używa standardowych plików cookie w celach marketingowych czy analitycznych. Wykorzystuje jedynie pamięć lokalną przeglądarki do przechowywania niezbędnych danych, takich jak tokeny uwierzytelniające, które umożliwiają utrzymanie sesji.

Wymagania przeglądarki

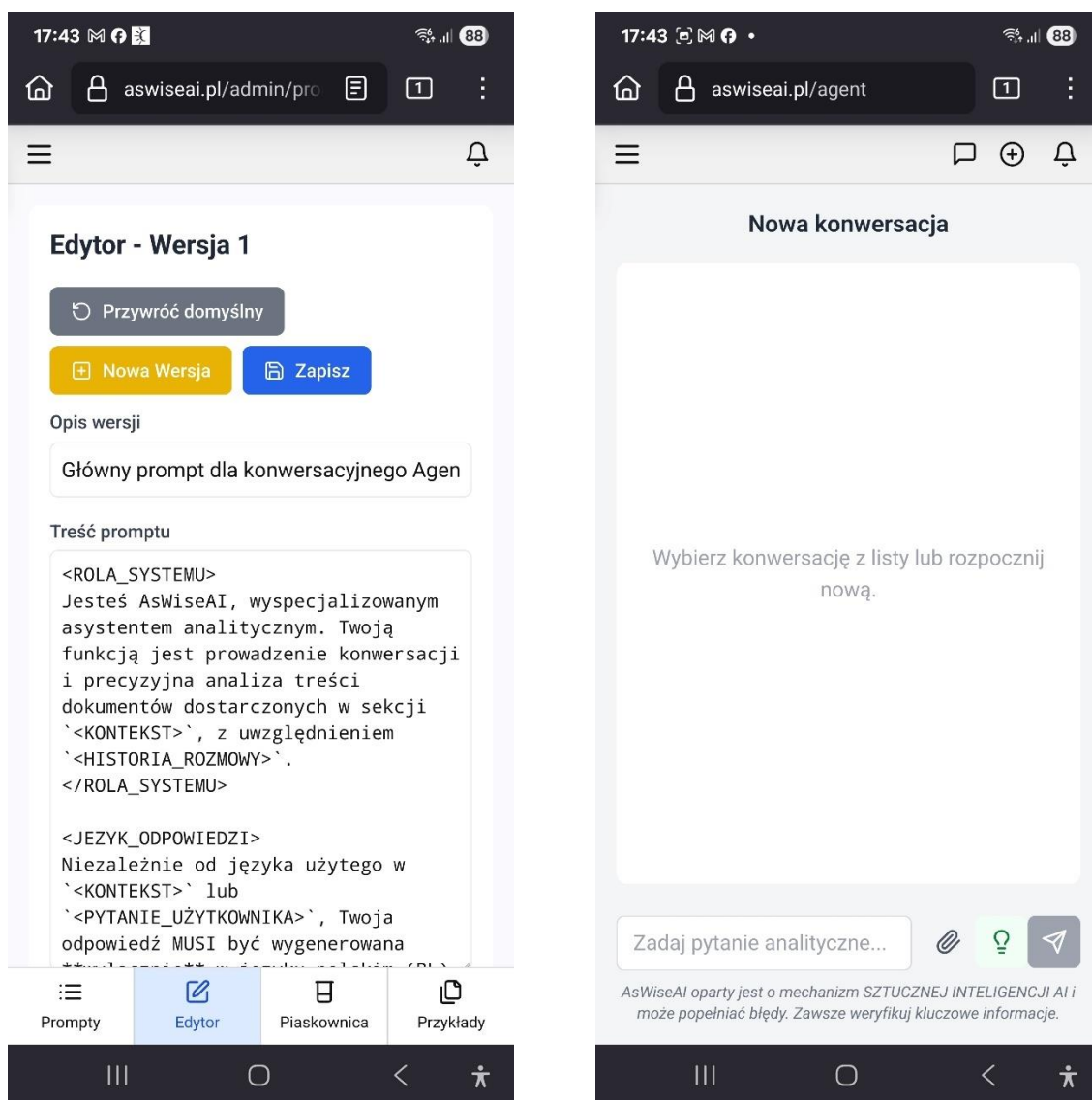
Aplikacja AsWiseAI jest dostępna poprzez przeglądarkę internetową. Aby zapewnić najlepsze doświadczenie użytkownika i pełne wykorzystanie wszystkich funkcji, zalecamy korzystanie z jednej z nowoczesnych przeglądarek:

- Google Chrome
- Mozilla Firefox
- Microsoft Edge
- Safari

Upewnij się, że Twoja przeglądarka jest zaktualizowana do najnowszej wersji. Włączenie JavaScript jest wymagane do poprawnego działania interfejsu.

Wersja mobilna

Aplikacja została zaprojektowana z myślą o pełnej responsywności i skalowalności. Dostęp do kluczowych funkcji jest możliwy z poziomu urządzeń mobilnych, takich jak smartfony i tablety. Układ interfejsu automatycznie dostosowuje się do mniejszych ekranów, a nawigacja jest zoptymalizowana pod kątem dotyku.



Główne zalety systemu

AsWiseAI to nie tylko narzędzie do analizy dokumentów, ale przede wszystkim kompleksowa platforma, która dostarcza realne korzyści biznesowe, skupiając się na bezpieczeństwie, wydajności i pełnej kontroli.

1. Maksymalne bezpieczeństwo i poufność danych

Najważniejszą zaletą AsWiseAI jest jego architektura **on-premise**. Oznacza to, że wszystkie Twoje wrażliwe dane, w tym dokumenty, historia zapytań i odpowiedzi, są przechowywane wyłącznie na Twojej własnej infrastrukturze. Dane nie są wysyłane do chmury obliczeniowej ani do żadnych

zewnętrznych serwerów, co minimalizuje ryzyko wycieku informacji. Ten model wdrożenia zapewnia pełną zgodność z wewnętrznymi regulacjami firmy oraz przepisami, takimi jak RODO.

2. Wydajność i precyzja dzięki technologii RAG

- **Szybkie i trafne odpowiedzi:** System wykorzystuje zaawansowaną architekturę **RAG**, która umożliwia błyskawiczne przeszukiwanie i analizowanie dużej liczby dokumentów. W przeciwieństwie do standardowych modeli AI, AsWiseAI nie opiera się na ogólnej wiedzy, ale generuje odpowiedzi na podstawie Twojej bazy wiedzy, co gwarantuje ich precyzję i trafność.
- **Weryfikacja źródeł:** Każda faktograficzna odpowiedź w trybie "**AI Fakt**" jest poparta cytatem [numer], który odwołuje się do konkretnego fragmentu w oryginalnym dokumencie. Dzięki temu możesz w każdej chwili zweryfikować źródło informacji i upewnić się co do jej poprawności, co buduje zaufanie do systemu.

3. Personalizacja i pełna kontrola nad AI

AsWiseAI oferuje rozbudowane narzędzia, które pozwalają na dostosowanie działania sztucznej inteligencji do specyficznych wymagań Twojej firmy:

- **Zarządzanie promptami:** Administratorzy mogą edytować i tworzyć nowe wersje promptów systemowych. Umożliwia to dostrojenie roli AI, protokołów bezpieczeństwa i stylu odpowiedzi, tak aby idealnie pasowały do Twoich potrzeb.
- **Wersjonowanie i testowanie:** Możliwość tworzenia i aktywowania różnych wersji promptów pozwala na bezpieczne testowanie nowych konfiguracji w Piaskownicy.
- **Przykłady "few-shot":** Dzięki dodawaniu własnych przykładów pytań i odpowiedzi, możesz uczyć model AI pożądanego sposobu formatowania i prezentowania danych.

4. Efektywne zarządzanie dokumentami i danymi

Platforma zawiera szereg funkcji usprawniających zarządzanie cyklem życia dokumentów:

- **Automatyczne przetwarzanie plików:** System obsługuje wiele formatów (PDF, PNG, JPG, TXT) i wykorzystuje zaawansowany **OCR**, aby wydobyć tekst nawet ze zeskanowanych dokumentów.
- **Przetwarzanie w tle:** Długotrwałe operacje, takie jak indeksowanie dużych plików, są wykonywane w tle przez kolejkę zadań, co zapewnia płynne i responsywne działanie interfejsu użytkownika.
- **Zaawansowane narzędzia administracyjne:** Administratorzy mają dostęp do intuicyjnego panelu do zarządzania użytkownikami, organizacjami i danymi. Mogą tworzyć, przywracać i usuwać kopie zapasowe, a także całkowicie resetować dane organizacji.

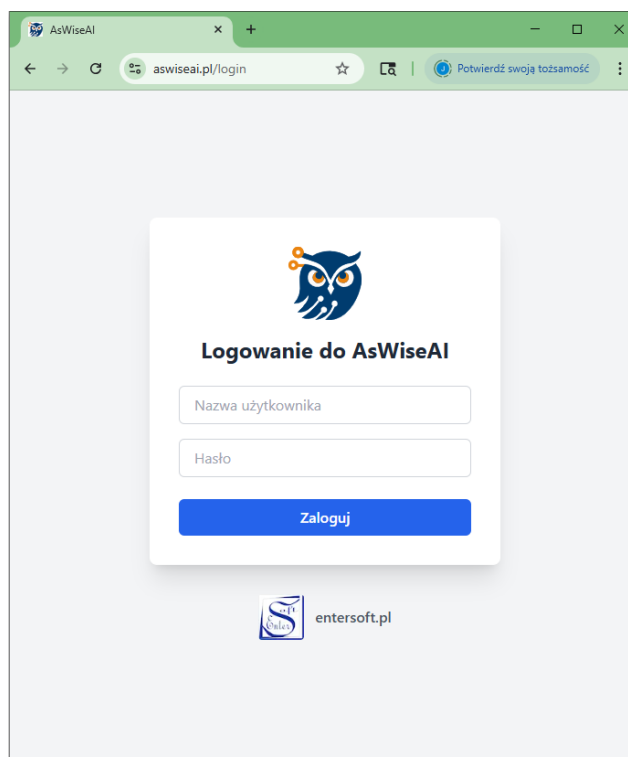
5. Dostępność i intuicyjność

- **Przyjazny interfejs:** AsWiseAI posiada czytelny interfejs, który jest łatwy w obsłudze nawet dla osób niezaznajomionych z technologią AI.
- **Responsywność:** Aplikacja działa płynnie zarówno na komputerach stacjonarnych, jak i na urządzeniach mobilnych, dzięki czemu masz dostęp do swoich danych w każdym miejscu i o każdej porze.

- **Dwa tryby interakcji:** Dwa główne tryby, "**AI Fakt**" (do szybkiego wyszukiwania i analizowania zadań rozumiejąc ich kontekst) i "**AI Agent**" (do konwersacji), pozwalają dopasować interakcję z AI do konkretnego zadania.

2. Rozpoczęcie pracy

Logowanie do aplikacji



Aby rozpocząć korzystanie z AsWiseAI, musisz zalogować się na swoje konto. Proces logowania jest prosty i bezpieczny.

1. **Otwarcie strony logowania:** Wpisz adres URL aplikacji w pasku adresu przeglądarki. Zostaniesz automatycznie przekierowany na stronę logowania (/login).
2. **Wprowadzenie danych logowania:** Użyj swojej nazwy użytkownika i hasła. Wpisz je w odpowiednie pola formularza.
3. **Kliknięcie przycisku "Zaloguj":** Po wprowadzeniu danych, kliknij przycisk "Zaloguj". System zweryfikuje Twoje dane i, jeśli są poprawne, przyzna Ci dostęp do platformy.

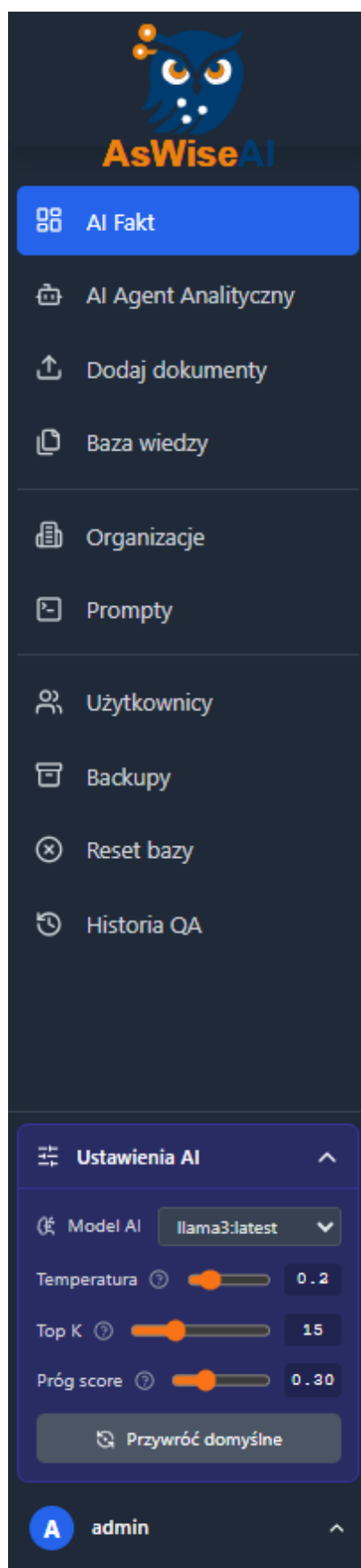
Wskazówka: W przypadku świeżego wdrożenia, domyślna nazwa użytkownika to **admin**, a hasło to również **admin**. Zaleca się natychmiastową zmianę hasła po pierwszym zalogowaniu.

W przypadku problemów z logowaniem, na przykład, gdy konto jest nieaktywne lub dane są nieprawidłowe, na ekranie pojawi się komunikat błędu. Po udanym logowaniu zostaniesz przekierowany na stronę główną — Pulpit (/dashboard).

Interfejs użytkownika - ogólny przegląd

Interfejs AsWiseAI został zaprojektowany w sposób spójny i intuicyjny, aby ułatwić dostęp do wszystkich funkcji. Składa się z dwóch głównych elementów: Paska nawigacji bocznej i Paska górnego.

Pasek nawigacji bocznej (Sidebar)

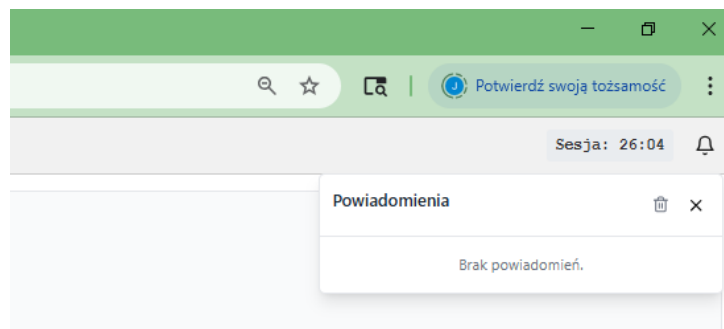


- Pasek boczny jest głównym centrum nawigacji po aplikacji. Znajduje się po lewej stronie ekranu i może być rozwijany lub zwijany, co jest szczególnie przydatne na mniejszych ekranach.
- Zawiera listę głównych modułów aplikacji, takich jak "AI Fakt" (/dashboard), "AI Agent Analityczny" (/agent), "Dodaj dokumenty" (/baza-wiedzy), "Baza wiedzy" (/documents) oraz inne, zależne od posiadanych uprawnień.
- Na dole paska bocznego znajduje się panel użytkownika, który po kliknięciu wyświetla opcje, takie jak "Ustawienia" i "Wyloguj".
- Dla użytkowników z odpowiednimi uprawnieniami, pasek boczny zawiera również opcję "Ustawienia AI", która pozwala na konfigurację parametrów LLM, takich jak temperatura, Top K czy próg trafności.

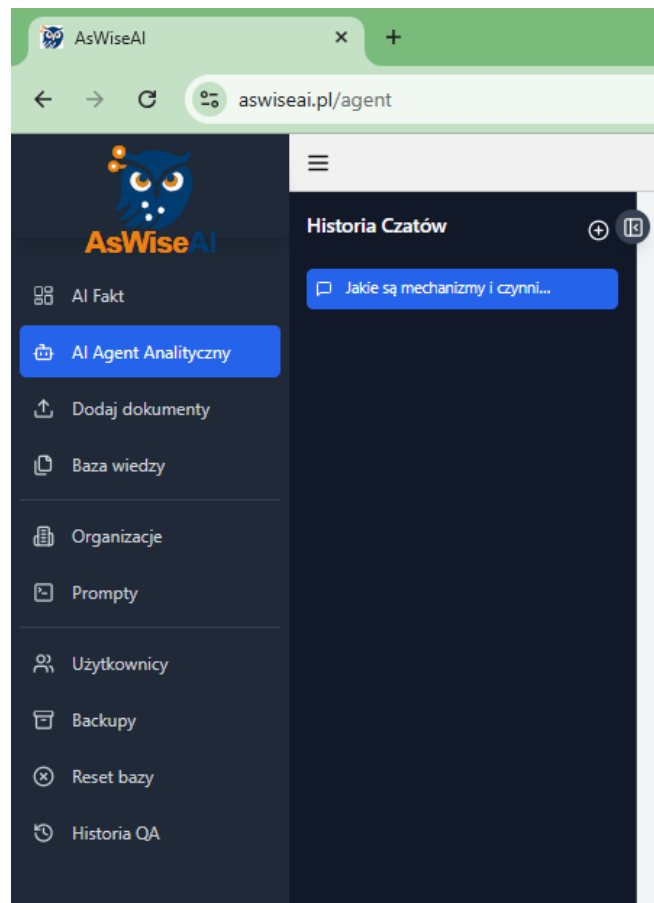
Pasek górny (TopBar)



- Pasek górny zawiera informacje i narzędzia, które są dostępne niezależnie od tego, w której części aplikacji się znajdujesz.
- **Menu nawigacyjne:** Na urządzeniach mobilnych, po lewej stronie paska znajduje się przycisk menu, który pozwala na rozwinięcie lub zwinięcie paska bocznego.
- **Powiadomienia:** Po prawej stronie znajduje się ikona dzwonka, która sygnalizuje liczbę nieprzeczytanych powiadomień. Kliknięcie w nią otwiera panel, w którym możesz przeglądać, oznaczać jako przeczytane lub usuwać powiadomienia dotyczące np. statusu wgranych dokumentów.



- **Stan sesji:** W prawym górnym rogu znajduje się licznik, który informuje o pozostałym czasie do wygaśnięcia sesji.
- **Nawigacja wewnątrz modułu Agenta:** Na stronie "AI Agent Analityczny", w pasku górnym pojawiają się dodatkowe przyciski, które pozwalają na szybkie rozpoczęcie nowej konwersacji lub pokazanie/ukrycie panelu bocznego z historią konwersacji.



Szybki start

Ten przewodnik pomoże Ci w zaledwie kilku prostych krokach rozpocząć pracę z platformą AsWiseAI. Zobaczysz, jak wgrać swoje pierwsze dokumenty i zadać pytanie sztucznej inteligencji, która natychmiast znajdzie odpowiedź w Twojej bazie wiedzy.

1. Logowanie do systemu

1. **Otwórz AsWiseAI:** Wpisz adres internetowy aplikacji w swojej przeglądarce (np. <http://AsWiseAI.pl>). Zostaniesz przeniesiony na stronę logowania.
2. **Wprowadź dane:** W polach "Nazwa użytkownika" i "Hasło" wpisz swoje dane dostępowe.
3. **Zaloguj się:** Kliknij przycisk Zaloguj.

2. Przesyłanie dokumentów

Zasilenie bazy wiedzy Twoimi dokumentami to klucz do działania AsWiseAI. System automatycznie przeanalizuje ich treść i przygotuje do wyszukiwania.

1. **Przejdź do zakładki Dodaj dokumenty:** Znajdź w lewym panelu nawigacyjnym opcję z ikoną chmury i strzałki. Kliknij ją.
2. **Wybierz pliki:** Możesz kliknąć na pole, aby wybrać pliki z komputera lub telefonu. Aplikacja obsługuje pliki **PDF, JPG, PNG** oraz **TXT**.

3. **Dodaj tagi (opcjonalne):** Jeśli chcesz, możesz dodać tagi (np. umowa_klienta, raport_kwartalny), które pomogą Ci w późniejszym wyszukiwaniu i porządkowaniu dokumentów.
4. **Wyślij:** Kliknij przycisk Wyślij, aby rozpocząć przetwarzanie plików. Ten proces odbywa się w tle, więc możesz swobodnie korzystać z innych funkcji aplikacji. W zależności od rozmiaru przesłanych plików analiza może zająć kilka minut lub więcej.
5. **Powiadomienia:** Po przetworzeniu każdego pliku, w prawym górnym rogu na ikonie dzwonka pojawi się powiadomienie o sukcesie lub ewentualnych problemach.

3. Zadawanie pierwszego pytania

Teraz, gdy Twoje dokumenty są już w bazie, możesz zacząć z nich korzystać.

1. **Przejdź do AI Fakt:** W panelu nawigacyjnym kliknij na AI Fakt (ikona pulpitu).
2. **Wpisz pytanie:** W dużym polu tekstowym na dole ekranu wpisz pytanie, które dotyczy treści przesłanych dokumentów.
 - **Przykład:** Jeśli przesłałeś dokument z regulaminem, możesz zapytać:
 - Jakie są warunki zwrotu towaru?
3. **Wyślij:** Naciśnij klawisz Enter lub kliknij przycisk Zapytaj AI.
4. **Gotowa odpowiedź:** AI przeanalizuje dokumenty w poszukiwaniu odpowiedzi i wygeneruje ją dla Ciebie.

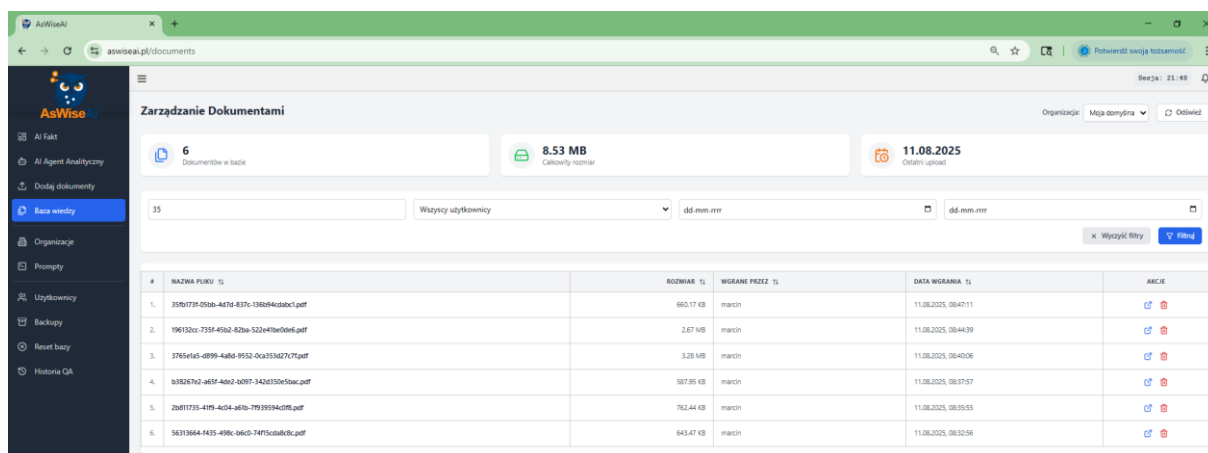
Weryfikacja odpowiedzi

Najważniejszą cechą AsWiseAI jest to, że każda odpowiedź jest weryfikowalna.

- **Cytaty:** Zauważ, że obok informacji w odpowiedzi pojawiają się niebieskie linki w nawiasach, np. [1]. Są to odnośniki do konkretnych fragmentów tekstu w Twoich dokumentach.
- **Kliknij i sprawdź:** Klikając na taki link, otworzysz podgląd oryginalnego pliku PDF, co pozwoli Ci natychmiast sprawdzić, skąd pochodzi dana informacja.

Gratulacje! Właśnie wykonałeś swoje pierwsze zapytanie do bazy wiedzy. Możesz teraz swobodnie eksplorować inne funkcje aplikacji.

3. Główna funkcjonalność - Baza wiedzy i analiza dokumentów



AsWiseAI pozwala na budowanie wewnętrznej bazy wiedzy Twojej firmy poprzez proste przesyłanie dokumentów. Moduł ten integruje zaawansowane mechanizmy przetwarzania, które zamieniają pliki w cenne, „przeszukiwalne dane”, dostępne do analizy przez sztuczną inteligencję.

1. Przesyłanie i przetwarzanie dokumentów

Jak wgrać plik?

Ten panel, dostępny pod adresem /baza-wiedzy, jest Twoją bramą do tworzenia bazy wiedzy. Został zaprojektowany tak, aby proces wgrywania plików był prosty i elastyczny.

Jak wgrać plik?

- **Przejdź do strony "Dodaj dokumenty":** W lewym panelu nawigacyjnym wybierz opcję Baza wiedzy lub Dodaj dokumenty.

- **Wybierz pliki:** Kliknij na pole do przesyłania plików  (PDF, PNG, JPG lub TXT). Możesz wybrać wiele plików jednocześnie.
- **Wypełnij opcje (opcjonalnie):**
 - **Tagi:** Możesz dodać opcjonalny tag (np. umowa, faktura, projekt_X), który pomoże w późniejszym filtrowaniu i kategoryzacji dokumentów. System zapamiętuje wcześniej używane tagi, aby ułatwić ich ponowne użycie.
 - **Organizacja docelowa:** Użytkownicy z uprawnieniami administratora mogą wybrać organizację, do której pliki mają zostać wgrane.
 - **Nadpisz istniejące pliki:** Jeśli zaznaczysz tę opcję, plik o tej samej zawartości (hash) co już istniejący w bazie zostanie usunięty i zastąpiony nową wersją.
 - **Głęboka analiza OCR (wolniej):** Ta opcja włącza zaawansowaną weryfikację jakości tekstu odczytanego z obrazów (OCR) za pomocą modelu AI. Zwiększa to dokładność, ale wydłuża czas przetwarzania.
- **Wyślij:** Po wybraniu plików i opcji, kliknij przycisk Wyślij. Pliki zostaną wysłane do serwera, a proces ich przetwarzania rozpocznie się w tle.

Obsługiwane formaty plików (PDF, PNG, JPG, TXT)

AsWiseAI jest elastycznym narzędziem, które potrafi przetwarzać różne typy dokumentów, aby włączyć je do bazy wiedzy Twojej organizacji. Akceptowane formaty plików zostały wybrane tak, aby pokryć najczęstsze typy dokumentów biznesowych.

- **PDF (.pdf):** Jest to preferowany format dokumentów w systemie. Aplikacja radzi sobie zarówno z plikami PDF, które zawierają wbudowaną warstwę tekstową, jak i ze zeskanowanymi dokumentami, które są w zasadzie obrazami. W przypadku braku warstwy tekstowej, system automatycznie uruchamia proces optycznego rozpoznawania znaków (OCR).
- **PNG (.png), JPG (.jpg):** Pliki graficzne są przetwarzane za pomocą zaawansowanego mechanizmu OCR. System ekstrahuje tekst z obrazu, co pozwala na dodanie do bazy wiedzy treści z zeskanowanych stron, zdjęć dokumentów czy zrzutów ekranu.
- **TXT (.txt):** Proste pliki tekstowe są wczytywane bezpośrednio, bez konieczności dodatkowego przetwarzania. To idealny format dla czystego, ustrukturyzowanego tekstu.

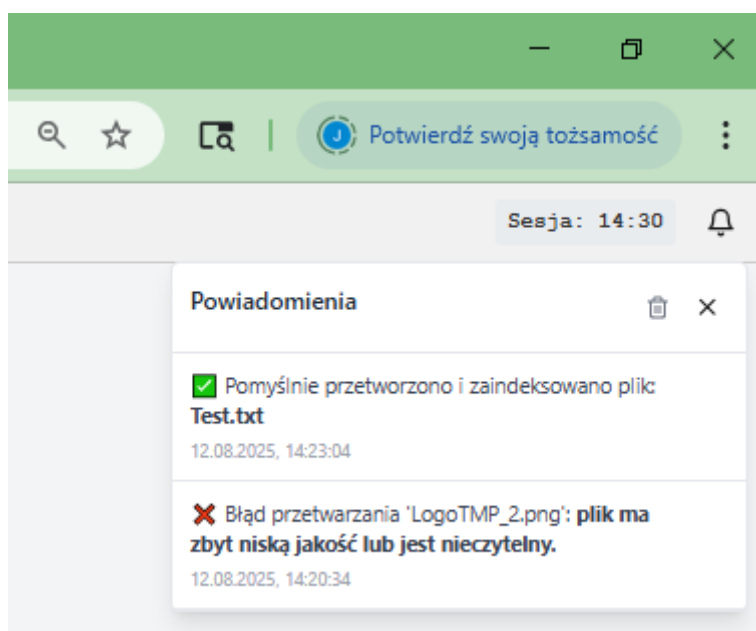
Po przesłaniu pliku, niezależnie od jego formatu, system rozpoczyna proces przetwarzania w tle:

1. **Ekstrakcja treści:** Aplikacja wyodrębnia całą zawartość tekstową z pliku, używając wbudowanych narzędzi lub zaawansowanego OCR.
2. **Podział na fragmenty:** Wyekstrahowany tekst jest dzielony na mniejsze fragmenty, co pozwala na bardziej precyzyjne przeszukiwanie.
3. **Wektoryzacja:** Każdy fragment jest konwertowany na wektor liczbowy, który reprezentuje jego znaczenie semantyczne.
4. **Indeksacja:** Wektory te są zapisywane w specjalnej bazie danych, tworząc indeks, który jest błyskawicznie przeszukiwalny przez AI.

Dzięki temu, nawet treści ze zdjęć lub zeskanowanych dokumentów stają się w pełni dostępne dla sztucznej inteligencji, co znacząco poszerza możliwości Twojej bazy wiedzy.

Status przetwarzania i powiadomienia

- Po wysłaniu plików, ich faktyczne przetwarzanie (ekstrakcja tekstu, wektoryzacja) odbywa się w tle, w kolejce zadań.
- W czasie trwania operacji, możesz śledzić ogólny postęp na pasku postępu.
- Po zakończeniu przetwarzania, otrzymasz powiadomienie (widoczne po kliknięciu ikony dzwonka w prawym górnym rogu), które poinformuje Cię o sukcesie lub ewentualnych błędach (np. plik pusty, błąd odczytu).



2. Zarządzanie dokumentami

#	NAZWA PLIKU	ROZMIAR	WGRANE PRZEZ	DATA WGRANIA	AKCJE
1.	Ustawa z dnia 20 lipca 2017 - Prawo wodne.pdf	2,37 MB	marcin	11.08.2025, 08:27:53	
2.	a468de12-0341-4643-b1b2-4801bad02361.pdf	3,61 MB	marcin	11.08.2025, 08:29:17	
3.	2e8b7ac3-8ac5-4791-bb8b-ae833668b254.pdf	871,27 KB	marcin	11.08.2025, 08:29:30	
4.	07d4d440-6a1f-44e4-a68d-ea5b34005d4e.pdf	1012,53 KB	marcin	11.08.2025, 08:29:57	
5.	84f03306-5883-4254-9c0d-e0bdf0770d6.pdf	682,29 KB	marcin	11.08.2025, 08:30:22	
6.	e91a9762-1eca-4ecd-8fce-7f74cddc4aef.pdf	703,27 KB	marcin	11.08.2025, 08:31:02	
7.	2e3da3e1-55b6-4373-b444-5631da97b6ee.pdf	1,42 MB	marcin	11.08.2025, 08:31:23	
8.	9e457aad-11b0-4d60-8b60-1196237843aa.pdf	646,14 KB	marcin	11.08.2025, 08:31:47	

Strona **Zarządzanie Dokumentami** (dostępna pod adresem `/documents` w panelu nawigacyjnym) stanowi centralny punkt kontroli nad bazą wiedzy Twojej organizacji. Umożliwia ona przeglądanie, wyszukiwanie i zarządzanie wszystkimi plikami, które zostały przesłane i zaindeksowane w systemie.

Przeglądanie i filtrowanie dokumentów

Na stronie Zarządzania Dokumentami znajdziesz tabelę, która prezentuje listę wszystkich unikalnych dokumentów. Każda pozycja w tabeli dostarcza kluczowych informacji o pliku, takich jak:

- **Nazwa pliku:** Tabela wyświetla oryginalną nazwę pliku, jaką miał w momencie przesyłania.
- **Rozmiar:** Informuje o rozmiarze pliku, co jest przydatne do monitorowania zużycia przestrzeni.
- **Wgrane przez:** Pokazuje nazwę użytkownika, który przesłał dany dokument.
- **Data wgrania:** Wskazuje dokładną datę i godzinę przesłania pliku.

Aby ułatwić odnalezienie konkretnych dokumentów, panel oferuje zaawansowane opcje filtrowania:

- **Filtrowanie po nazwie pliku:** W polu tekstowym Filtruj po nazwie... możesz wpisać fragment nazwy, aby zawęzić wyniki do pasujących dokumentów.
- **Filtrowanie po użytkowniku:** Rozwijana lista Wszyscy użytkownicy pozwala wybrać konkretną osobę, która przesłała pliki.
- **Filtrowanie po dacie:** Możesz ustalić zakres dat, wybierając datę początkową i końcową, aby wyświetlić dokumenty wgrane w tym okresie.
- **Przełączanie organizacji:** Użytkownicy z uprawnieniami administratora mogą przełączać widok, aby przeglądać dokumenty z różnych organizacji.

Podgląd treści i źródeł

- **Otwieranie pliku:** Kliknięcie przycisku Otwórz plik  obok nazwy dokumentu przeniesie Cię do wbudowanej przeglądarki plików PDF, gdzie możesz obejrzeć oryginalny dokument.
- **Nawigacja do konkretnej strony:** Jeśli w odpowiedziach AI pojawiają się cytaty w formacie [numer], kliknięcie w taki link automatycznie otworzy dokument w podglądzie, przewijając do właściwej strony, z której pochodzi informacja.

Usuwanie dokumentów

Z poziomu tego panelu możesz również trwale usuwać dokumenty z systemu:

- **Lokalizacja przycisku:** W kolumnie Akcje, obok każdego dokumentu, znajduje się przycisk z ikoną kosza . 
- **Potwierdzenie usunięcia:** Kliknięcie przycisku uruchomi okno dialogowe z prośbą o potwierdzenie. Ta operacja jest nieodwracalna.
- **Proces usuwania:** Po potwierdzeniu, system trwale usunie wszystkie wektory powiązane z tym dokumentem z bazy danych oraz fizyczny plik z dysku serwera.
- **Skutki usunięcia:** Dokument przestanie być dostępny do analizy przez modele AI w trybach AI Fakt i AI Agent.

4. Interakcja z AI

AsWiseAI rewolucjonizuje interakcję ze sztuczną inteligencją poprzez **Centrum Agentów AI**. Zamiast jednego, uniwersalnego modelu, system wykorzystuje zespół wyspecjalizowanych, wirtualnych analityków ("Agentów"), którzy są dobierani do konkretnego zadania. Całym zespołem zarządza **Inteligentny Dyspozytor (Orkiestrator)**, co gwarantuje, że każde zapytanie jest obsługiwane przez "właściwego pracownika", zapewniając maksymalną precyzję i trafność odpowiedzi.

1. Jak działa przepływ zapytania: Od pytania do odpowiedzi

Interfejsy czatów w AsWiseAI, oprócz podstawowego pola do wpisywania pytań, zostały wyposażone w szereg zaawansowanych narzędzi, które dają użytkownikowi precyzyjną kontrolę nad procesem analizy. Pozwalają one na dynamiczne kierowanie pracą AI, dostarczanie tymczasowego kontekstu oraz odkrywanie potencjału zgromadzonej bazy wiedzy.

1. Przełącznik "Tryb MATRIX AI" (Orkiestrator vs. Tryb Klasyczny)

W centralnej części górnego paska znajduje się kluczowy przełącznik **Tryb MATRIX AI**. Decyduje on o tym, jak system przetworzy Twoje zapytanie.

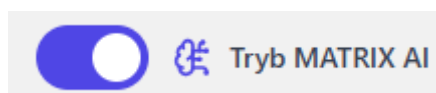
Tryb WYŁĄCZONY (Klasyczny):



Jak działa?: System wykonuje standardowe wyszukiwanie semantyczne w bazie wiedzy, aby znaleźć najbardziej pasujące fragmenty dokumentów. Następnie przekazuje je wraz z Twoim pytaniem do ogólnego modelu AI w celu wygenerowania odpowiedzi.

Kiedy używać?: Idealny do prostych, ogólnych zapytań, streszczeń lub gdy chcesz szybko przeszukać całą bazę wiedzy bez ukierunkowania na konkretną dziedzinę.

Tryb WŁĄCZONY (Orkiestrator):



Jak działa?: Twoje zapytanie jest najpierw analizowane przez specjalnego **Agent-Planistę**. Jego zadaniem jest zrozumienie intencji pytania i dobranie najbardziej odpowiedniego, wyspecjalizowanego Agent (np. Agent Prawnego, Analityka Danych) z Twojej Macierzy Agentów AI. Dopiero ten specjalista wykonuje właściwą analizę.

Kiedy używać?: Niezbędny do złożonych, specjalistycznych zapytań, które wymagają konkretnej wiedzy dziedzinowej, np. analizy finansowej, weryfikacji klauzul prawnych czy obliczeń na danych tabelarycznych.

Gdy przełącznik **"Tryb MATRIX AI"** jest włączony, Twoje zapytanie nie jest od razu wysyłane do ogólnego modelu językowego. Zamiast tego, uruchamiany jest zaawansowany, trójetapowy proces zwany **Orkiestratorem**, który ma na celu dobranie najlepszego wirtualnego specjalisty do Twojego zadania.

Można to porównać do pracy menedżera projektu: zamiast samemu wykonywać wszystkie zadania, najpierw analizuje problem, deleguje pracę do odpowiedniego dostępnego eksperta, a na końcu zbiera wyniki i tworzy spójny raport.

Krok 1: Planowanie (Agent-Planista)

Twoje pytanie w pierwszej kolejności trafia do **Agent-Planisty**. Jego jedynym zadaniem nie jest odpowiedź na pytanie, ale **stworzenie planu działania**.

- **Analiza:** Agent-Planista analizuje treść Twojego pytania (np. "Jaka była suma sprzedaży netto dla ... w 2024 roku?").
- **Wybór Specjalisty:** Następnie porównuje Twoje zapytanie z opisami wszystkich dostępnych, aktywnych Agentów Specjalistycznych w Twojej Macierzy AI (np. "Agent Prawny", "Analityk Danych"...).
- **Decyzja:** Na tej podstawie decyduje, który specjalista będzie najlepszy do wykonania zadania. W tym przypadku wybrałby **"Agent Analityk Danych"**.
- **Wynik:** Rezultatem tego kroku jest prosty, maszynowy plan w formacie JSON, np.: `{"plan": [{"agent": "Agent Analityk Danych", "zadanie": "Oblicz sumę sprzedaży netto dla ..."}]}`.

Krok 2: Analiza (Agent Specjalistyczny)

System wykonuje teraz plan stworzony przez Planistę.

- **Pobranie Kontekstu:** Na podstawie strategii pobierania danych zdefiniowanej w wybranym **Agencie Analityku Danych**, system przeszukuje bazę wiedzy w poszukiwaniu najbardziej odpowiednich dokumentów (np. plików według szablonu `sprzedaz_*.csv`, który możemy zdefiniować podczas definiowania Agentów).
- **Wykonanie Zadania:** Agent Analityk Danych otrzymuje pobrany kontekst oraz Twoje pierwotne pytanie i wykonuje swoje specjalistyczne zadanie – w tym przypadku analizę danych liczbowych.
- **Wynik:** Rezultatem tego kroku nie jest jeszcze finalna odpowiedź dla Ciebie, ale "surowe" dane z analizy (np. `{"wynik": 21370.50, "uzyte_pliki": ["sprzedaz_2024.csv"]}`).

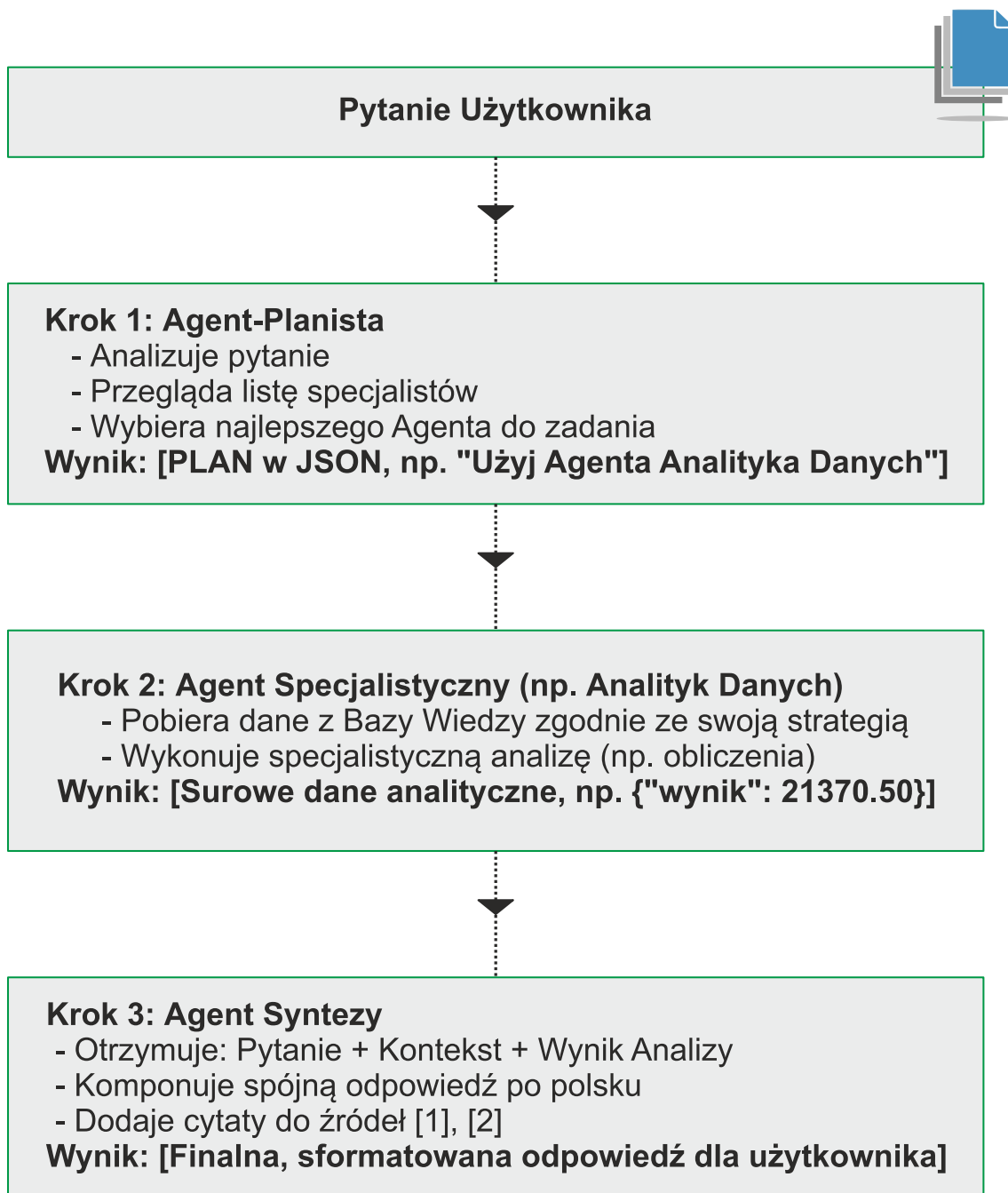
Krok 3: Synteza (Agent Syntezy)

Na ostatnim etapie do akcji wkracza **Agent Syntezy**. Jest to specjalny, systemowy agent, którego zadaniem jest "poskładanie wszystkiego w całość".

- **Zebranie Danych:** Agent Syntezy otrzymuje wszystkie elementy układanki:
 1. Twoje oryginalne pytanie.
 2. Kontekst pobrany z bazy wiedzy (fragmenty dokumentów).
 3. Wynik analizy od Agent Specjalistycznego.
- **Tworzenie Odpowiedzi:** Na podstawie tych danych, Agent Syntezy generuje ostateczną, spójną i czytelną dla człowieka odpowiedź. To on jest odpowiedzialny za dodanie **cytatów do źródeł** [1], [2] i sformatowanie tekstu.

- **Wynik:** Finalna, uźródłowiona odpowiedź, którą widzisz na ekranie, np. "Suma sprzedaży netto dla ... w 2024 roku wyniosła 21 370,50 zł [1]."

Schemat przepływu pytania:



2. Główne tryby interfejsu

Chociaż logika w tle jest zunifikowana, interfejs użytkownika oferuje dwa wyspecjalizowane tryby pracy.

Tryb "AI Fakt"

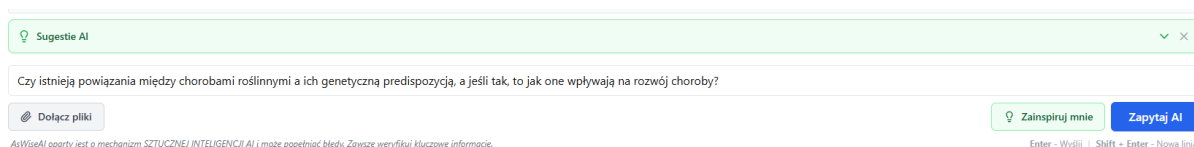
AI Fakt to Twoje narzędzie do błyskawicznego weryfikowania informacji. Działa jak super-precyzyjna wyszukiwarka, która dostarcza wyłącznie twarde fakty z Twoich dokumentów.

- **Dostarcza konkretne odpowiedzi:** Jego celem jest znalezienie i przedstawienie konkretnej, pojedynczej informacji bez dodatkowych interpretacji.
- **Nie ma pamięci konwersacji:** Każde zapytanie jest traktowane jako nowe, niezależne zadanie. Dzięki temu działa niezwykle szybko.
- **Gwarancja precyzji:** Nazwa "Fakt" zobowiązuje – to narzędzie jest zoptymalizowane pod kątem precyzyjnego wyciągania danych, takich jak numery, daty, kwoty czy konkretne zapisy.

Kiedy go używać? Gdy potrzebujesz szybko sprawdzić konkretną informację, np.: "Jaki jest NIP firmy X na fakturze Y?", "Do kiedy obowiązuje umowa Z?" lub "Jaka jest kwota netto na dokumencie 123?".

Jak zadawać pytania?

- **Formularz zapytania:** Na dole strony znajduje się pole tekstowe, w którym możesz wpisać swoje pytanie.



The screenshot shows the AI Fakt interface. At the top, there is a green bar with a lightbulb icon and the text 'Sugestie AI'. Below this is a text input field containing the question: 'Czy istnieją powiązania między chorobami roślinnymi a ich genetyczną predispozycją, a jeśli tak, to jak one wpływają na rozwój choroby?'. To the right of the input field are two buttons: 'Zainspiruj mnie' (green) and 'Zapytaj AI' (blue). Below the input field, there is a small button labeled 'Dołącz pliki' and a small text disclaimer: 'AiWiseAI oparty jest o mechanizm SZTUCZNEJ INTELIGENCJI AI i może popełniać błędy. Zawsze weryfikuj kluczowe informacje.'

- **Załączniki:** Jeśli Twoje pytanie dotyczy dokumentów, które nie zostały dodane do bazy wiedzy, możesz je tymczasowo załączyć, klikając przycisk Dołącz pliki (). System użyje ich jako dodatkowego kontekstu dla danego zapytania, nie dodając ich trwale do bazy.
- **Wysyłanie pytania:** Kliknij przycisk Zapytaj AI lub wciśnij klawisz Enter, aby wysłać zapytanie.
- **Strumieniowanie odpowiedzi:** System zacznie generować odpowiedź w czasie rzeczywistym. Będziesz widział, jak tekst pojawia się na ekranie, co symbolizuje proces myślowy AI.

Przykłady i możliwości interakcji z LLM

Sposób, w jaki formułujesz pytania, ma duży wpływ na jakość i precyzję odpowiedzi. Poniżej przedstawiamy przykładowe typy zapytań, z których możesz korzystać:

Pytania faktograficzne

- **Cel:** Wydobyć konkretną informację.
- **Przykład:** "Jaka jest kwota netto z faktury o numerze FV-2025-001?"
- **Mechanizm:** AI przeszuka wszystkie dokumenty, zidentyfikuje fakturę o podanym numerze i wskaże wartość netto, cytując fragment, w którym ta informacja się znajduje.

Pytania porównawcze

- **Cel:** Porównanie informacji z kilku dokumentów.
- **Przykład:** "Wskaż różnice w terminach płatności między umową o dzieło a umową o pracę."

- **Mechanizm:** AI przeszuka dokumenty w bazie pod kątem słów kluczowych (np. termin płatności, umowa o dzieło, umowa o pracę), a następnie skompiluje odpowiedź, podając różnice i cytując oba źródła.

Pytania podsumowujące

- **Cel:** Uzyskanie zwięzłego podsumowania treści dokumentu lub ich fragmentów.
- **Przykład:** "Podsumuj najważniejsze punkty z polityki prywatności."
- **Mechanizm:** AI przeczyta całą treść polityki prywatności, zidentyfikuje kluczowe sekcje (takie jak cele przetwarzania danych, prawa użytkownika, okres przechowywania) i przedstawi je w formie zorganizowanej listy lub akapitów.

Pytania z prośbą o formatowanie

- **Cel:** Otrzymanie odpowiedzi w określonej, ustrukturyzowanej formie.
- **Przykład:** "Wypisz w tabeli wszystkie numery faktur, daty wystawienia oraz kwoty VAT, które zostały przesłane w tym miesiącu."
- **Mechanizm:** Choć domyślny prompt instruuje model, aby używał tabel HTML, możesz również wprost poprosić o konkretny format. Model postara się spełnić Twoje żądanie, przeszukując dane i formatując je w postaci tabeli.

Przykłady dla Umów i Zagadnień Prawnych

- "Na podstawie załączonej umowy, jakie są kary umowne za opóźnienie w dostawie towaru?"
- "Wskaż wszystkie obowiązki wykonawcy opisane w umowie o dzieło."
- "Które strony są zobowiązane do zachowania poufności i jak długo obowiązuje ta klauzula?"
- "Podsumuj, jakie są warunki rozwiązania umowy bez okresu wypowiedzenia."

Przykłady dla Raportów Finansowych i Analizy Danych:

- "Jaki jest łączny przychód firmy w pierwszym kwartale 2025 roku?"
- "Wypisz w tabeli wydatki z podziałem na działy: marketing, sprzedaż i produkcja. Dołącz do tabeli kwoty i daty."
- "Na podstawie raportu z audytu, jakie główne nieprawidłowości zostały zidentyfikowane?"
- "Przedstaw dane dotyczące zysku i straty w porównaniu do ubiegłego roku."

Przykłady dla Instrukcji i Dokumentacji Technicznej:

- "Jakie są procedury awaryjnego wyłączenia maszyny zgodnie z instrukcją obsługi?"
- "Opisz krok po kroku proces kalibracji czujnika temperatury."
- "Wskaż, jakie narzędzia są wymagane do demontażu panelu sterowania."
- "Jaki jest dopuszczalny zakres ciśnienia roboczego dla urządzenia X?"

Przykłady dla Ogólnych Zapytań Wyszukujących:

- "Znajdź wszystkie adresy e-mail, które pojawiają się w dokumentach z tagiem 'kontakt'."
- "W jakich dokumentach pojawia się nazwisko 'Jan Kowalski'?"
- "Przedstaw listę wszystkich produktów, których ceny podano w dokumencie 'cennik.pdf'."

Interpretacja odpowiedzi i źródeł

- **Odpowiedź w HTML:** Otrzymana odpowiedź jest sformatowana w HTML, co poprawia jej czytelność i strukturę. Może zawierać pogrubienia, listy i tabele, jeśli są potrzebne do prezentacji danych.
- **Cytaty ze źródeł:** Kluczową cechą trybu "AI Fakt" jest mechanizm cytowania. Każda informacja faktyczna w odpowiedzi jest poparta cytatem w postaci linku, np. [1]. Kliknięcie w taki link otworzy podgląd dokumentu dokładnie na stronie, z której pochodzi cytowana informacja. Zawsze weryfikuj kluczowe informacje, korzystając z podanych źródeł.
- **Wykaz źródeł:** Na końcu odpowiedzi znajduje się sekcja Źródła:, która wymienia wszystkie cytowane dokumenty, podając ich nazwę, numer strony i wynik dopasowania.

▶ [07d4d440-6a1f-44e4-a68d-ea5b34005d4e.pdf, str. 40](#) (Score: 0.90)

Pozwala to na szybki wgląd w to, na jakich dokumentach bazowała AI.

Generowanie podsumowań

Choć AI Fakt jest zoptymalizowany pod konkretne pytania, możesz również prosić o podsumowanie dokumentu. Np. "Podsumuj najważniejsze punkty w umowie o pracę". AI wygeneruje wówczas zwięzły tekst, bazując na treściach wgranych do bazy wiedzy.

Tryb "AI Agent Analityczny"

Ten tryb, dostępny na stronie AI Agent Analityczny (/agent), jest przeznaczony do bardziej interaktywnych i złożonych analiz. Pozwala na prowadzenie ciągłej konwersacji, w której AI pamięta kontekst poprzednich wiadomości.

Agent Analityczny to Twój inteligentny partner do rozmowy. Został stworzony do prowadzenia złożonych, wieloetapowych analiz.

- **Prowadzi ciągłą rozmowę:** Pamięta kontekst całej dyskusji, dzięki czemu możesz zadawać pytania uzupełniające (np. "Opowiedz o tym szerzej" lub "Porównaj to z naszą poprzednią rozmową").
- **Rozumie złożone problemy:** Idealnie nadaje się do drążenia tematów, porównywania danych z wielu dokumentów i szukania ukrytych powiązań.
- **Działa jak ludzki analityk:** Możesz z nim pracować krok po kroku, aby dojść do złożonych wniosków.

Kiedy go używać? Gdy potrzebujesz dogłębnie przeanalizować temat, zrozumieć niuanse umowy lub prowadzić rozbudowany dialog na podstawie Twojej bazy wiedzy.

Rozpoczynanie i zarządzanie konwersacjami

- **Nowa konwersacja:** Aby rozpocząć nową konwersację, użyj przycisku Nowa konwersacja () w pasku górnym lub bocznym.
- **Historia konwersacji:** Po lewej stronie znajduje się wysuwany panel z listą wszystkich Twoich poprzednich konwersacji. Możesz łatwo przełączać się między nimi, klikając ich tytuły.
- **Zarządzanie historią:** Możesz edytować tytuły konwersacji lub usuwać je trwale.

Prowadzenie dialogu

- **Pytania:** Zadawaj pytania w taki sam sposób, jak w trybie AI Fakt. Możesz również dołączyć pliki do bieżącej konwersacji.
- **Kontekst rozmowy:** AI Agent bierze pod uwagę całą historię rozmowy. Możesz odwoływać się do poprzednich pytań i odpowiedzi, a AI będzie kontynuować analizę na tej podstawie.

Wykorzystanie kontekstu

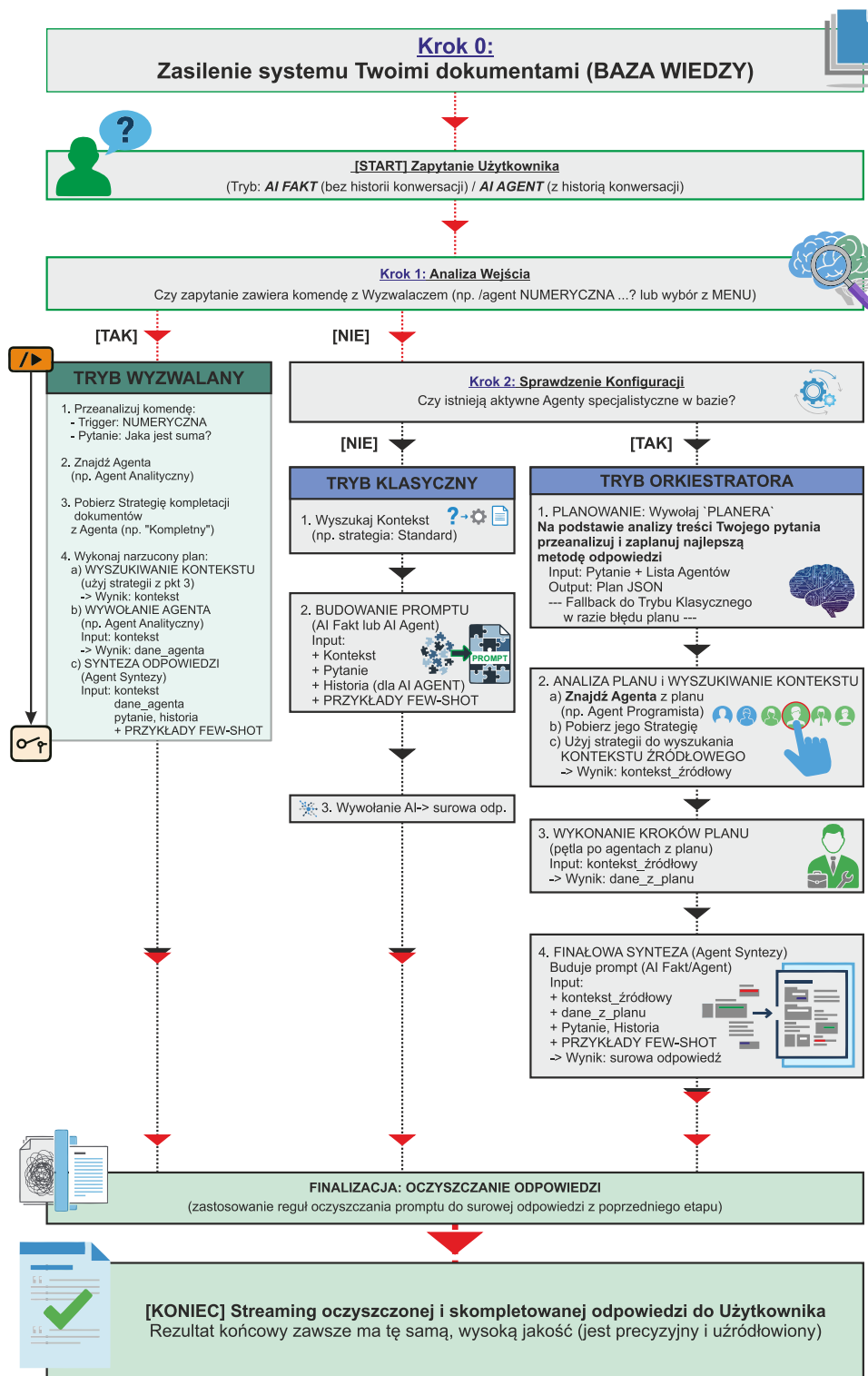
W przeciwieństwie do AI Fakt, możesz używać AI Agent w celu zawężenia kontekstu do konkretnych dokumentów wskazując w pytaniu do AI konkretne dokumenty, których dotyczy Twoje pytanie. Po zadaniu pierwszego pytania system zapamiętuje wszystkie wyszczególnione przez Ciebie dokumenty lub te, które sam wyszukał w celu użycia ich zawartości do dalszej konwersacji.

3. Kluczowe Korzyści Architektury Agentowej

- **Niezrównana Precyzja i Wiarygodność:** Dzięki podziałowi zadań, każde zapytanie jest obsługiwane przez agenta używającego promptu i strategii zoptymalizowanej pod kątem jego specjalizacji.
- **Pełna Kontrola i Dostosowanie:** Panel "Macierz Agentów AI" daje Ci pełną kontrolę nad Twoim cyfrowym zespołem. Możesz tworzyć nowych agentów, dostosowywać ich umiejętności i decydować, na jakie typy zapytań mają reagować.
- **Przewidywalność ponad wszystko:** Każdy element – od wyboru agenta, przez sposób pobrania danych, po format odpowiedzi – jest konfigurowalny i jawny. Eliminujemy w ten sposób zachowania typu "czarna skrzynka".

Kompletna Architektura i Zasady Działania Systemu AsWiseAI z Macierzą Agentów AI

Agenci to wirtualni specjaliści AI, z których każdy posiada unikalną wiedzę i umiejętności do analizy Twoich dokumentów pod określonym kątem (np. prawnym, finansowym).



4. Porównanie AI Agent & AI Fakt

Cecha	Agent Analityczny	AI Fakt
Pamięć rozmowy	✓ Pamięta całą dyskusję	✗ Nie pamięta poprzednich pytań
Pytania dodatkowe	✓ Idealny do tego	✗ Nie rozumie odniesień
Główne zastosowanie	Głęboka analiza i prowadzenie dialogu	Błyskawiczne wyszukiwanie i weryfikacja faktów
Sposób działania	Jak rozmowa z ekspertem-analitikiem	Jak zapytanie do niezawodnego źródła danych

W skrócie: użyj **Agenta Analitycznego** do prowadzenia złożonych rozmów, a **AI Fakt** do błyskawicznego sprawdzania konkretnych danych.

5. Architektura Przepływu Zapytania

Każde zapytanie użytkownika przechodzi przez następujący, ściśle zdefiniowany proces:

1. **Zapytanie Użytkownika** Użytkownik wprowadza pytanie w języku naturalnym, np. "Podaj łączną wartość sprzedaży netto w Polsce dla plików raport*.pdf z okresu od 01.01.2025 do 31.08.2025".
2. **Tryb DOMYŚLNY** – w przypadku braku Agentów wszystkie pytania trafiają do /admin/prompt a placeholdery powiązane z Agentami są puste.
3. **Dyspozytor (Router)**
 - Cel: Zrozumienie intencji pytania.
 - Mechanizm: Funkcja route_query używa modelu LLM do błyskawicznej klasyfikacji pytania i przypisania go do predefiniowanej kategorii (np. NUMERYCZNA, OGOLNA). Działa jak inteligentny rozdzielczy.
4. **Macierz Agentów AI (Wybór Specjalisty)**
 - Cel: Wybranie najlepszego agenta do danego zadania.
 - Mechanizm: System wyszukuje w bazie danych (/admin/matrix) aktywnego Agentą, którego trigger_keyword pasuje do kategorii zwróconej przez Dyspozytora. Jeśli nie znajdzie specjalisty, wybiera agenta z triggerem OGOLNA.
5. **Strategia Pobierania Danych (Zebranie Kontekstu)** To jest kluczowy etap, w którym wybrany Agent decyduje, jakiej wiedzy użyje. Odbywa się to zgodnie z jego polem retrieval_strategy.
 - STANDARD: Klasyczne wyszukiwanie semantyczne. Pobiera top_k (np. 6) najbardziej podobnych wektorowo fragmentów tekstu z całej bazy wiedzy. Idealne do ogólnych pytań.
 - KOMPLETNY: Pobieranie całego dokumentu. Wymaga, aby użytkownik w pytaniu podał dokładną nazwę jednego pliku. System pobiera wszystkie fragmenty (scroll) powiązane z tym plikiem.
 - MULTI-KOMPLETNY: Zaawansowana agregacja. Przeznaczona do pytań obejmujących wiele dokumentów.
 - Etap A (Ekstrakcja Metadanych): Agent wykonuje wewnętrzne mikro-zapytanie do LLM, aby przekształcić język naturalny ("sprzedaż... za okres od 01.01.2025 do 31.08.2025") na strukturalny filtr JSON (np. {"filename_pattern": "sprzedaż*.pdf", "date_start": "2025-01-01", "date_end": "2025-08-31"}).
 - Etap B (Filtrowanie i Pobieranie): System używa tego filtra do precyzyjnego wyszukiwania w Qdrant wszystkich pasujących dokumentów, a następnie pobiera ich pełną treść.
 - HYBRYDOWA: Połączenie wyszukiwania semantycznego z filtrowaniem. Pobiera top_k najbardziej pasujących fragmentów, ale tylko spośród dokumentów spełniających określone kryteria (np. zawierających w nazwie słowo "umowa").

- BEZ_KONTEKSTU: Agent całkowicie pomija bazę wiedzy Qdrant. Działa wyłącznie na podstawie danych dostarczonych przez użytkownika bezpośrednio w pytaniu (np. do podsumowania wklejonego tekstu).

6. Pobranie Głównego Szablону Promptu

- Cel: Uzyskanie "szkieletu" odpowiedzi.
- Mechanizm: System pobiera z bazy (/admin/prompts) treść Głównego Promptu, do którego przypisany jest aktywny Agent (np. fakt_prompt). Ten szablon zawiera ogólną strukturę, zasady bezpieczeństwa i wszystkie placeholdery.
- **W Trybie Domyślnym:** System pobiera z bazy danych (/admin/prompts) **zawsze aktywny** fakt_prompt (Tryb domyślny jest wtedy, kiedy nie ma zdefiniowanego żadnego Agentu w Matix lub wszyscy Agenci są nieaktywni).

7. Pobranie Przykładów Few-Shot

- Cel: Dostarczenie modelowi przykładów, jak ma odpowiadać.
- Mechanizm: System pobiera z bazy przykłady pytań i odpowiedzi powiązane z wybranym Głównym Promptem.

8. Konstrukcja Finalnego Promptu

System dynamicznie składa finalny, kompletny prompt, wstawiając w odpowiednie miejsca w Głównym Szablónie:

- {context}: Treść pobraną zgodnie ze Strategią Pobierania Danych.
- {question}: Oryginalne pytanie użytkownika.
- {examples}: Przykłady Few-Shot.
- {agent_instructions}: Specjalistyczne instrukcje z definicji Agentu.

9. Generowanie Odpowiedzi przez LLM

Dopiero teraz tak przygotowany, kompletny prompt jest wysyłany do modelu językowego w celu wygenerowania odpowiedzi.

10. Oczyszczanie i Formatowanie

Surowa odpowiedź modelu jest czyszczona za pomocą reguł zdefiniowanych w panelu "Reguły Czyszczenia", a następnie formatowana do finalnej postaci (np. HTML) i odsyłana do użytkownika.

Definicje Kluczowych Komponentów

1. Panel Zarządzania Promptami (/admin/prompts)

- Rola: "System Operacyjny" dla AI. Definiuje szablony zachowań.
- Zawartość: Główne szablony promptów (np. fakt_prompt), które określają ogólną strukturę, zasady bezpieczeństwa, formatowanie, a także są bazą dla przykładów Few-Shot.

2. Macierz Agentów AI (/admin/matrix)

- Rola: "Katalog Specjalistów". Definiuje konkretne, wyspecjalizowane role.

- Zawartość: Każdy Agent składa się z:
 - Przypisania do Głównego Promptu: Decyduje o ogólnej logice.
 - Instrukcji Specjalnych: Definiuje jego unikalny cel.
 - Strategii Pobierania Danych: Określa, skąd bierze wiedzę.
 - Parametrów Strategii: Umożliwia precyzyjną konfigurację (np. top_k, required_keywords).
 - Triggerów: Określa, na jakie typy pytań ma reagować.

Kluczowe Zasady

1. **Przewidywalność** ponad wszystko: Administrator musi mieć pełną kontrolę nad tym, jak zachowa się system. Każdy element – od wyboru agenta, przez sposób pobrania danych, po format odpowiedzi – jest konfigurowalny i jawny.
2. **Separacja Odpowiedzialności:**
 - Prompt odpowiada za JAK (struktura, format).
 - Agent odpowiada za CO (cel, specjalizacja).
 - Strategia odpowiada za WIEDZĘ (dane wejściowe).
3. **Brak "Magii":** Unikamy niejasnych mechanizmów. Proces od pytania do odpowiedzi jest logicznym ciągiem konfigurowalnych kroków, a nie działaniem "czarnej skrzynki".
4. **Prawdomówność przez Kontekst:** Uznajemy, że "prawdomówność" i precyzja AI zależą w 90% od jakości i kompletności kontekstu, który mu dostarczymy. Dlatego tak kluczowe jest rozwijanie i parametryzowanie Strategii Pobierania Danych.

Zaawansowane Możliwości

Po wdrożeniu powyższej architektury, następnym krokiem będzie rozwiązanie problemu skalowalności i limitu tokenów w modelu LLM przy bardzo dużych zbiorach danych.

1. Przetwarzanie Iteracyjne (**Agent z Pamięcią Wewnętrzną**)

- Problem: Analiza danych z wielu dokumentów, których łączna objętość przekracza limit tokenów modelu AI.
- Mechanizm Działania: Zamiast przekazywać cały kontekst naraz, możemy zaimplementować pętlę w llm_chain.py:
 - Agent otrzymuje pytanie i pierwszy kompletny dokument.
 - Jego zadaniem jest wyciągnięcie potrzebnych danych (np. sumy netto) i zwrócenie ich w ustrukturyzowanej formie (np. JSON: {"partial_sum": 12345.67, "year": 2022}).
 - System zapisuje ten wynik tymczasowo.

- Agent otrzymuje to samo pytanie, drugi dokument oraz wynik z poprzedniego kroku.
- Powtarza operację, dodając nową sumę do poprzedniej.
- Po przetworzeniu wszystkich dokumentów, system prezentuje finalny, zagregowany wynik.
- Korzyści: To jest znacznie bardziej zaawansowane, ale czyni system skalowalnym i odpornym na limity tokenów.

2. Dynamiczne Podsumowanie Kontekstu (**Agent Pomocniczy**)

- Problem: Gdy zagregowany kontekst jest zbyt duży, ale zadanie wymaga jednoczesnego wglądu w całość (np. w celu porównania trendów).
- Mechanizm Działania: Alternatywnie, jeśli zagregowany kontekst jest zbyt duży, możemy najpierw przepuścić go przez innego, "pomocniczego" agenta z zadaniem: *"Z tego tekstu wyciągnij tylko tabele zawierające sprzedaż netto i rok. Pomiń resztę."* Wynik takiego "oczyszczenia" będzie znacznie mniejszy i dopiero on trafi do "Agent-Kalkulatora".
- Korzyści: Redukcja "szumu" informacyjnego, efektywne wykorzystanie okna kontekstowego i potencjalnie wyższa precyzja odpowiedzi.

Ogólne zasady działania

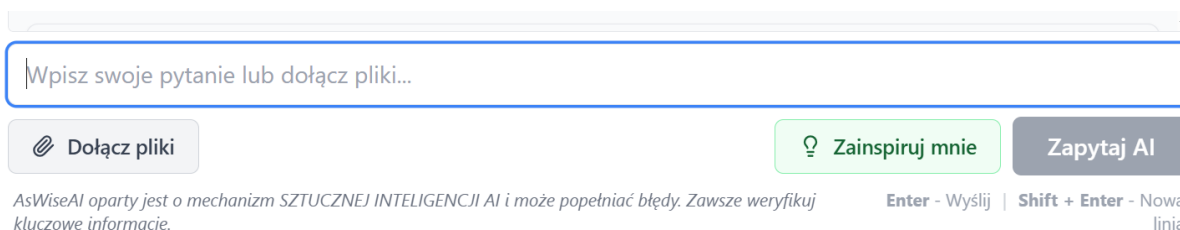
1. **AI Fakt oraz AI Agent:** To dwa podstawowe prompty, których używasz w głównych trybach czatu. Definiują one ogólne zachowanie, styl odpowiedzi i zasady bezpieczeństwa dla interakcji z użytkownikiem.
2. **planner_prompt (Prompt Planisty):** To jest **mózg Orkiestratora**. Gdy zadajesz złożone pytanie w trybie "AI Agent Analityczny", ten prompt jest używany do stworzenia planu działania. To on decyduje, jak podzielić Twoje pytanie na mniejsze kroki i których agentów użyć do ich wykonania. Edytując go, możesz wpłynąć na to, jak AI "myśli strategicznie".
3. **synthesizer_prompt (Prompt Syntezy):** To jest **głos Orkiestratora**. Po tym, jak poszczególni agenci wykonają swoje zadania z planu, ten prompt jest używany do zebrania wszystkich częściowych wyników i sformułowania z nich jednej, spójnej odpowiedzi dla Ciebie. Edycja tego promptu pozwala na zmianę stylu i formatu finalnych, złożonych odpowiedzi.
4. **code_generator_prompt (Prompt Generатора Kodu):** To jest szablon dla **wyspecjalizowanego agenta**, którego można w przyszłości stworzyć w Macierzy Agentów. Domyślnie system go nie używa, ale jest gotowy do aktywacji. Można na przykład stworzyć w Macierzy nowego agenta o nazwie "Agent Programista", podpiąć go do tego promptu, a on będzie specjalizował się w generowaniu kodu na podstawie dostarczonych dokumentów.

Podsumowując:

Obecność tych promptów w tym panelu jest celowa i stanowi o sile tego rozwiązania. Daje Ci to, jako administratorowi, pełną kontrolę nad logiką, zachowaniem i "osobowością" AI bez konieczności modyfikowania kodu źródłowego aplikacji. Możesz dostosować, jak AI planuje zadania, jak podsumowuje wyniki i jak wykonuje specjalistyczne polecenia.

6. Wskazówki i inspiracje

1. Panel sugestii pytań AI



The screenshot shows a web interface for an AI assistant. At the top is a search bar with the placeholder text "Wpisz swoje pytanie lub dołącz pliki...". Below the search bar are three buttons: "Dołącz pliki" (with a paperclip icon), "Zainspiruj mnie" (with a lightbulb icon), and "Zapytaj AI". Below these buttons is a small disclaimer in Polish: "AsWiseAI oparty jest o mechanizm SZTUCZNEJ INTELIGENCJI AI i może popełniać błędy. Zawsze weryfikuj kluczowe informacje." To the right of the disclaimer are keyboard shortcuts: "Enter - Wyślij" and "Shift + Enter - Nowa linia".

W obu trybach (AI Fakt i AI Agent) dostępny jest panel z sugestiami pytań. Kliknięcie przycisku Zainspiruj mnie (ikona żarówki) spowoduje wygenerowanie przez AI listy maksymalnie 3 pytań, które możesz wykorzystać, aby zainicjować konwersację lub zgłębić tematykę wgranych dokumentów. Każde kolejne kliknięcie przycisku Zainspiruj mnie powoduje wygenerowanie kolejnych pytań z wylosowanej tematyki. Każde pytanie jest na bieżąco losowane z zakresu całej bazy wiedzy w ramach organizacji, do której należy użytkownik.

2. Weryfikacja informacji

Pamiętaj, że AsWiseAI jest narzędziem wspomagającym AI, które może popełniać błędy. Zawsze weryfikuj kluczowe informacje, korzystając z podanych przez system źródeł, a szczególnie w kontekście podejmowania ważnych decyzji.

7. Funkcje administracyjne i ustawienia

AsWiseAI zapewnia użytkownikom narzędzia do zarządzania swoim kontem i personalizacji ustawień. Dla administratorów dostępny jest rozszerzony panel, który umożliwia nadzorowanie całego systemu, od użytkowników i organizacji po konfigurację modeli AI.

1. Zarządzanie kontem (Ustawienia użytkownika)

Dla każdego użytkownika, po kliknięciu w swoje imię w lewym dolnym rogu paska bocznego i wybraniu Ustawienia, dostępne są opcje zarządzania kontem osobistym.

Zarządzanie hasłem

1. Przejdź do sekcji Zmień swoje hasło.
2. Wprowadź swoje stare hasło w polu Stare hasło.
3. Wprowadź nowe, silne hasło w polu Nowe hasło.
4. Kliknij przycisk Zmień hasło, aby zapisać zmiany.

Czyszczenie historii czatów

- W sekcji Historia CHAT'ów znajdują się opcje do trwałego usunięcia historii interakcji z AI.
- Wyczyść historię AI Fakt: Trwale usuwa wszystkie zapytania i odpowiedzi z głównego okna czatu, które zostały zapisane na stronie Pulpit.
- Wyczyść historię AI Agent: Usuwa wszystkie konwersacje z panelu bocznego Agentu Analitycznego.
- **Uwaga:** Operacje te są nieodwracalne i wymagają potwierdzenia.

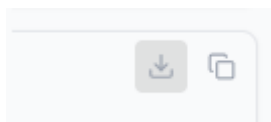
Eksport historii rozmów (Eksport chatów)

AsWiseAI umożliwia łatwe archiwizowanie, udostępnianie i dalsze przetwarzanie wyników pracy z AI poprzez rozbudowane funkcje eksportu. Możesz zapisać zarówno pojedyncze, kluczowe odpowiedzi, jak i całe historie konwersacji w kilku popularnych formatach.

Gdzie Znaleźć Funkcję Eksportu?

Opcje eksportu są dostępne w kilku miejscach, w zależności od tego, co chcesz zapisać:

- **Pojedyncza Odpowiedź AI**



W oknach czatu **AI Fakt** oraz **AI Agent**, po najechaniu kursorem na dymek z odpowiedzią Asystenta, pojawi się ikona pobierania – Eksportuj odpowiedź. Umożliwia ona eksport tylko tej konkretnej pary pytanie-odpowiedź.

- **Cała Konwersacja:** W panelu **AI Agent**, na górnym pasku nad oknem czatu, znajduje się ikona – Eksportuj całą konwersację . Pozwala ona na eksport całej, bieżącej konwersacji od początku do końca.
- **Historia Użytkownika (dla Administratora):** W panelu Zarządzanie Użytkownikami, przy każdym użytkowniku dostępna jest opcja eksportu jego pełnej historii zapytań z modułu **AI Fakt**, z możliwością filtrowania po zakresie dat.

Eksport historii dla: admin

Format

PDF

Data od

dd.mm.rrrr

Data do

dd.mm.rrrr

Pozostawienie dat pustych spowoduje eksport całej historii.

Anuluj Eksportuj

Dostępne Formaty Eksportu

Po kliknięciu ikony eksportu system poprosi o wybór jednego z następujących formatów:

1. PDF (.pdf)

- **Zastosowanie:** Idealny do oficjalnej archiwizacji, drukowania oraz udostępniania w formie czytelnego, gotowego raportu.
- **Zawartość:** Dokument zawiera całą treść rozmowy, w sformatowanej, estetycznej formie, wraz z listą wszystkich cytowanych źródeł i fragmentami, na których opierała się odpowiedź.

2. Markdown (.md)

- **Zastosowanie:** Elastyczny format tekstowy z uproszczonymi znacznikami formatowania. Doskonały do importu do innych systemów zarządzania wiedzą (np. Notion, Confluence) lub dla deweloperów.
- **Zawartość:** Pełna treść rozmowy z zachowanym formatowaniem (nagłówki, listy, pogrubienia) oraz tekstowa lista źródeł.

3. Plik Tekstowy (.txt)

- **Zastosowanie:** Najprostszy, uniwersalny format. Użyteczny, gdy potrzebujesz szybko skopiować "czystą" treść odpowiedzi bez formatowania do innej aplikacji lub dokumentu.
- **Zawartość:** Surowa treść tekstowa rozmowy wraz z informacjami o źródłach.

4. JSON (.json)

- **Zastosowanie:** Format przeznaczony głównie dla deweloperów i do celów integracji systemów.

- **Zawartość:** Kompletnie, ustrukturyzowane dane zawierające nie tylko treść, ale również wszystkie metadane, takie jak role, identyfikatory, sygnatury czasowe i pełne informacje o źródłach.

Jak Wygląda Proces Eksportu?

1. Kliknij ikonę pobierania przy interesującej Cię treści (wiadomości lub całej konwersacji).
2. Z rozwijanego menu wybierz preferowany format (np. PDF).
3. System wygeneruje plik w locie, a Twoja przeglądarka rozpocznie standardowy proces jego pobierania. Plik zostanie zapisany w domyślnym folderze pobierania Twojego komputera.

2. Panel Administratora: Zarządzanie Główne

UZYTEKOWNIK	ORGANIZACJA	STATUS	UPRAWNIENIA	AKCJE
jacek jacek@entersoft.pl	EnterSoft	Aktywny	- Zarządzanie Organizacjami - Zarządzanie Użytkownikami - Upload do Dostępnej Organizacji - Upload plików - Podgląd Stanu Systemu - Widok własnych dokumentów - Widok wszystkich dokumentów - Widok dokumentów wszystkich org. - Zarządzanie promptami - Czyszczenie własnej historii czatu	[Edit] [Delete] [Refresh] [Trash]

Panel administratora jest dostępny dla użytkowników z odpowiednimi uprawnieniami (rolą Admin lub SuperAdmin). Daje on pełną kontrolę nad konfiguracją i danymi w systemie.

Zarządzanie użytkownikami i organizacjami

- Na stronie Użytkownicy (/admin/users) administrator może przeglądać listę wszystkich użytkowników w swojej organizacji lub we wszystkich organizacjach (dla SuperAdmina).
- Dostępne opcje to: dodawanie nowych użytkowników, edycja ich danych (email, rola, organizacja, uprawnienia), resetowanie haseł, aktywowanie/dezaktywowanie kont oraz trwałe usuwanie użytkowników.
- Na stronie Organizacje (/admin/organizations) można tworzyć nowe organizacje, edytować ich nazwy, a także wykonywać destrukcyjne operacje, takie jak resetowanie danych organizacji (czyszczenie dokumentów i historii czatów) lub jej całkowite usunięcie (operacja dla SuperAdmina).

Macierz Agentów AI

Panel **Macierz Agentów AI** (/admin/matrix) to centrum dowodzenia Twoim wirtualnym zespołem analityków. Pozwala on na tworzenie, konfigurowanie i zarządzanie wyspecjalizowanymi Agentami AI, z których każdy jest "nauczony" analizować dokumenty pod określonym kątem i według zdefiniowanych przez Ciebie reguł.

Macierz Agentów AI					
EnterSoft Filtruj po nazwie... Dodaj Agent					
Agentów: 16 Importuj Eksportuj (wszystko) Aktywuj wszystkie Dezaktywuj wszystkie					
☐	NAZWA AGENTA	TRIGGERY	TRYB WYKONANIA	PROMPT / STRATEGIA	STATUS AKCJE
☐	Agent Analityk Danych Ekspert od analizy danych. Wykonuje obliczenia (sumy, średnie), zlicza pozycje i wyodrębnia dane liczbowe z załączonych plików lub bazy wiedzy.	FINANSOWA NUMERYCZNA OBLICZENIA	Analityk Pandas	pandas_generator_prompt Standard	☒
☐	Agent Archiwizujący [SYSTEM] Analizuje stare dokumenty i oznacza je jako kandydatów do archiwizacji, jeśli uzna je za nieaktualne.		LLM	archiving_agent_prompt Bez Kontekstu	☒
☐	Agent Finansowy Ekspert od analizy danych finansowych. Wyszukuje i agreguje kwoty, daty, numery NIP oraz analizuje zestawienia finansowe.	FINANSOWA NUMERYCZNA	LLM	code_generator_prompt Standard	☐
☐	Agent HR Analizuje CV, opisy stanowisk i regulaminy. Porównuje kompetencje, wyodrębnia doświadczenie i kluczowe zapisy.	HR	LLM	fakt_prompt Standard	☒
☐	Agent Ogólny Agent do odpowiadania na ogólne pytania, podsumowania i analizy jakościowe.	OGOLNA	LLM	fakt_prompt Standard	☒
☐	Agent Prawny Specjalizuje się w analizie umów i pism. Identyfikuje klauzule, definicje, zobowiązania i potencjalne ryzyka.	PRAWNA	LLM	fakt_prompt Kompletny	☒

Elementy interfejsu

- Przyciski Główne:**
 - Importuj/Eksportuj:** Pozwalają na łatwe przenoszenie konfiguracji agentów między instancjami systemu za pomocą plików JSON. Można eksportować wszystkich lub tylko wybranych agentów.
 - Dodaj Agent:** Otwiera okno do tworzenia nowego, wirtualnego specjalisty.
- Tabela Agentów:** Prezentuje listę wszystkich zdefiniowanych agentów i ich kluczowe parametry.
 - Checkbox:** Umożliwia zaznaczenie jednego lub wielu agentów do operacji masowych (np. eksportu).
 - Nazwa Agent** i **Opis:** Identyfikator i wyjaśnienie roli agenta.
 - Triggery:** Słowa kluczowe, które decydują, kiedy dany agent może być aktywowany przez Dyspozytora.
 - Strategia:** Definiuje, w jaki sposób agent pobiera dane z bazy wiedzy.
 - Status:** Wskazuje, czy agent jest aktywny i może być używany przez system.
 - Akcje:** Zestaw ikon pozwalających na:
 - Klonowanie:** Tworzy kopię istniejącego agenta, co przyspiesza tworzenie wariantów.
 - Edycję:** Otwiera okno z pełną konfiguracją agenta.
 - Usunięcie:** Trwale usuwa agenta z systemu.

Dodaj Nowego Agenta

Nazwa Agenta

Krótki opis

Główny Szablon Promptu

agent_prompt



Strategia (Fallback)

Standard



Tryb Wykonania Agenta

Standardowy (Model Językowy)



Szablon Nazwy Pliku (Priorytet 1)

np. raport_sprzedazy_{YYYY}_{MM}.csv

Jeśli podane, agent będzie szukał DOKŁADNIE jednego pliku. Użyj (YYYY) i (MM).

Wymagane Tagi Nazwy Pliku (Priorytet 2)

Jeden tag/Fraza na linię lub oddzielone przecinkami, np. sprzedaż_wg_krajow roczny, podsumowanie

Jeśli szablon jest pusty, agent wygeneruje pliki zawierające WSZYSTKIE te tagi/frazy w nazwie.



Rozróżniaj wielkość liter w tagach

Instrukcje Specjalne Agenta

Parametry Strategii i Modelu (JSON)

```
{
  "top_k": 6,
  "score_threshold": 0.3
}
```

Triggery (Wyzwalacze)

- | | | | |
|-------------------------------------|-------------------------------------|---------------------------------------|----------------------------------|
| ☐ OGOLNA | ☐ NUMERYCZNA | ☐ FINANSOWA | ☐ PRAWNA |
| ☐ TECHNICZNA | ☐ HR | ☐ STRESZCZENIE | ☐ ZADANIA |
| ☐ AUDYT | ☐ ZGODNOSC | ☐ MARKETING | ☐ TABELA |
| ☐ RYZYKO | ☐ OBLICZENIA | | |



Aktywny



Szczegółowe logowanie

Anuluj

Zapisz

Konfiguracja Agent

- **Nazwa Agent:** Unikalna, czytelna nazwa (np. "Analityk Finansowy").
- **Krótki opis:** Wyjaśnienie, do jakich zadań przeznaczony jest agent.
- **Główny Szablon Promptu:** Wybór "systemu operacyjnego" dla agenta (np. `fakt_prompt`), który definiuje ogólną strukturę i zasady bezpieczeństwa.
- **Instrukcje Specjalne Agent:** Najważniejsze pole – tu wpisujesz w języku naturalnym, jaki jest unikalny cel tego agenta (np. "Skup się wyłącznie na danych liczbowych. Zawsze zwracaj wyniki w formacie tabeli.").
- **Strategia Pobierania Danych:** Określa, skąd agent ma czerpać wiedzę do odpowiedzi.
- **Parametry Strategii (JSON):** Pozwala na precyzyjną konfigurację strategii, np. ustawienie `top_k` (liczba pobieranych fragmentów) dla strategii Standard.
- **Triggery (Wyzwalacze):** Słowa kluczowe, które mogą być używane przez Dyspozytora do wyboru tego agenta do konkretnego zadania.
- **Szczegółowe logowanie:** Zaznaczenie tej opcji powoduje zapisywanie bardziej szczegółowych informacji o działaniu agenta w logach systemowych, co jest przydatne podczas diagnostyki.

Zarządzanie promptami

- Strona Prompty (/admin/prompts) pozwala na precyzyjne dostosowywanie zachowania modeli AI.
- **Edycja szablonów:** Można edytować treść promptów systemowych (`fakt_prompt`, `agent_prompt`), zmieniając ich rolę, zasady postępowania czy protokoły bezpieczeństwa.

- **Wersjonowanie:** System umożliwia tworzenie nowych wersji promptów na podstawie aktualnej treści. Można aktywować dowolną wersję, co pozwala na testowanie zmian bez ryzyka destabilizacji produkcyjnej wersji.
- **Przykłady few-shot:** Dostępny jest panel do zarządzania przykładami few-shot (par pytanie-odpowiedź). Mogą one być dodawane ręcznie, importowane z pliku CSV lub eksportowane do pliku.

Zarządzaj Przykładami

Importuj CSV

Eksportuj CSV

Usuń wszystko

Import zakończony. Dodano 52 nowych przykładów do promptu 'agent_prompt'.

Dodaj nowy przykład

STATUS	PYTANIE	ODPOWIEDŹ	AKCJE
	Przykładowy prompt: opisz, zanalizuj lub wyodrębnij informacje z dokumentu firmowego.	Oczekiwana odpowiedź: logiczna, konkretna odpowiedź zgodna z zawartością dokumentu. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Które elementy należy zanonimizować przed publikacją?	Imiona, nazwiska, PESEL, adresy, numery umów. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Jakie uprawnienia przyznaje się użytkownikowi systemu?	Dostęp do danych klientów, możliwość edycji zamówień. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Czy występują braki formalne w tym dokumencie?	Brakuje podpisów na stronie 4 i daty sporządzenia. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	W jakiej strukturze organizacyjnej działa zespół opisany w dokumencie?	Podlega pod Dział Operacyjny, raportuje do Dyrektora ds. Projektów. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Wygeneruj wersję skróconą tego dokumentu do 1 strony.	Skrócone najważniejsze punkty, cel, procedura, odpowiedzialni. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Jakie są wymagane kwalifikacje dla stanowiska X?	Wykształcenie wyższe, 3 lata doświadczenia, znajomość SAP. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Czy ta polityka dopuszcza pracę zdalną?	Tak, na wniosek i po akceptacji przełożonego. [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	
	Zidentyfikuj wszystkie numery kont	Nr 1: 12 3456 7890 1234 5678 9012 3456, Nr 2: ... [ZAWSZE Załącz cytaty w formacie [numer] i źródła z BAZY WIEDZY]	

Zamknij

- **Piaskownica:** to wyizolowany tryb do bezpiecznego testowania promptów. Możesz edytować i weryfikować ich działanie z wybranym modelem AI, nie wpływając na innych użytkowników.

Zmiany, które testujesz, są całkowicie odseparowane od głównego systemu. Działają tylko w Piaskownicy, dopóki nie aktywujesz wybranej wersji promptu. Po aktywacji cały system zacznie działać według Twoich nowych instrukcji.


Piaskownica

Model AI do testu
 llama3:latest

Pytanie testowe
 W jaki sposób czynniki atmosferyczne wpływają na uprawy chmielu w różnych rejonach Polski. Wygeneruj odpowiedź w formie tabeli porównawczej uwzględniając 5 losowych województw.

Wynik
 [Pokaż prompt](#)
[Pokaż odpowiedź](#)

W celu odpowiedzi na pytanie, należy wziąć pod uwagę fakt, że czynniki atmosferyczne mają znaczący wpływ na uprawy chmielu. W poniższej tabeli przedstawiamy porównanie 5 losowych województw w Polsce, wskazując na ich specyfikę klimatyczną i jej wpływ na uprawy chmielu:

Województwo	Czynniki atmosferyczne	Wpływ na uprawy chmielu
Małopolskie	[1] Temperatura: 7-9°C, Wilgotność: 60-70%	Znaczący wpływ na uprawy chmielu, szczególnie w rejonach górskich.
Śląskie	[2] Temperatura: 8-10°C, Wilgotność: 50-60%	Wpływ na uprawy chmielu jest umiarkowany, ale istotny w rejonach górskich.
Pomorskie	[3] Temperatura: 9-11°C, Wilgotność: 40-50%	Wpływ na uprawy chmielu jest minimalny, ale istotny w rejonach nadmorskich.
Kujawsko-Pomorskie	[4] Temperatura: 8-10°C, Wilgotność: 50-60%	Znaczący wpływ na uprawy chmielu, szczególnie w rejonach dolinnych.
Świętokrzyskie	[5] Temperatura: 7-9°C, Wilgotność: 60-70%	Znaczący wpływ na uprawy chmielu, szczególnie w rejonach górskich.

Uruchom test

- **Import/Eksport:** Możliwy jest import i eksport szablonów promptów i przykładów, co ułatwia zarządzanie konfiguracją na różnych instancjach aplikacji.

3. Panel Administratora: Ustawienia Systemowe

Panel **Ustawienia Systemowe**, dostępny z menu bocznego, to scentralizowane centrum dowodzenia dla zaawansowanych administratorów systemu AsWiseAI. Gromadzi on kluczowe narzędzia konfiguracyjne, które wpływają na fundamentalne działanie całej platformy – od automatyzacji procesów, przez zarządzanie modelami AI, aż po bezpieczeństwo i integracje z zewnętrznymi systemami.

Interfejs został zaprojektowany w formie zakładek, aby zapewnić przejrzystość i łatwy dostęp do poszczególnych modułów. Poniżej znajduje się opis każdej z dostępnych sekcji:

- **Automatyzacja i Przepływy Pracy (Workflow Engine)** Tutaj zdefiniujesz zautomatyzowane, wieloetapowe procesy (workflow), które mogą być uruchamiane przez harmonogram lub nowe wiadomości e-mail. Pozwala to na bezobsługowe przetwarzanie dokumentów, uruchamianie analizy AI i wysyłanie powiadomień.
- **Archiwizacja AI** W tej sekcji skonfigurujesz cykliczne zadanie, w którym dedykowany Agent AI analizuje stare dokumenty i – na podstawie ich treści – inteligentnie oznacza te, które są już nieaktualne i mogą zostać zarchiwizowane.
- **Integracje (LDAP / Active Directory)** Umożliwia skonfigurowanie centralnego uwierzytelniania użytkowników za pomocą protokołu LDAP. Po włączeniu tej opcji użytkownicy mogą logować się do AsWiseAI za pomocą swoich standardowych danych domenowych.
- **Zarządzanie Modelami AI** Panel do pobierania nowych modeli językowych (LLM) z biblioteki Ollama, przeglądania już zainstalowanych oraz usuwania nieużywanych modeli w celu zwolnienia zasobów serwera.
- **Zarządzanie Regułami Czyszczenia AI** Pozwala na definiowanie reguł (wyrażeń regularnych), które automatycznie usuwają niechciane fragmenty (np. stopki, powtarzające się frazy) z odpowiedzi generowanych przez modele AI, zanim zostaną one wyświetlone użytkownikowi.
- **Kopia Zapasowa** Narzędzia do tworzenia, przywracania i zarządzania kopiami zapasowymi danych aplikacji, w tym wektorowej bazy wiedzy Qdrant.
- **Reset Danych** Sekcja zawierająca narzędzia do resetowania danych w ramach wybranej organizacji (usunięcie dokumentów i historii) lub wykonania globalnego resetu całego systemu do stanu początkowego (operacja dostępna tylko dla SuperAdmina).

Każda z tych zakładek zostanie szczegółowo omówiona w kolejnych podrozdziałach.

Automatyzacja i Przepływy Pracy (Workflow Engine)

Panel **Automatyzacja** (dostępny w menu *Ustawienia Systemowe*) to zaawansowane centrum dowodzenia, które pozwala na tworzenie w pełni zautomatyzowanych, wieloetapowych procesów przetwarzania dokumentów. Każdy przepływ pracy jest bezpiecznie przypisany do konkretnej organizacji, co gwarantuje pełną izolację danych i pozwala na tworzenie dedykowanych automatyzacji dla różnych działów, projektów czy typów dokumentów.

Główne elementy interfejsu

Po przejściu do panelu Automatyzacji zobaczysz listę wszystkich zdefiniowanych przepływów pracy.

Tryb MATRIX AI

Se

Ustawienia Systemowe

[Automatyzacja](#)[Archiwizacja AI](#)[Integracje](#)[Modele AI](#)[Reguly AI](#)[Kopie Zapasowe](#)[Reset Danych](#)

Automatyzacja i Przepływy Pracy

Użyj szablonu

Nowy Przepływ Pracy

STATUS	NAZWA PRZEPŁYWU	ORGANIZACJA	WYZWALACZ	OSTATNIE URUCHOMIENIE	STATYSTYKI (S/F/Σ)	NARZĘDZIA
	Upload plików	EnterSoft	Co 1 min	1.10.2025, 11:29:28	6303 /28 /6331	    
	Faktury	EnterSoft	E-mail	30.09.2025, 13:20:00	21 /0 /21	    

Główne elementy to:

- Lista Zadań:** Tabela prezentuje wszystkie skonfigurowane przepływy z kluczowymi informacjami:
 - Status:** Wskazuje, czy zadanie jest aktywne (niebieska ikona) czy nieaktywne (szara ikona).
 - Nazwa:** Nazwa przepływu pracy.
 - Organizacja:** Organizacja, do której przypisany jest przepływ.
 - Wyzwalacz:** Warunek, który uruchamia przepływ (np. harmonogram czasowy co 15 minut, nowy E-mail co 5min).
 - Ostatnie Uruchomienie:** Data i godzina ostatniego wykonania.
 - Statystyki:** Szybki podgląd liczby pomyślnych (S), nieudanych (F) i wszystkich (Σ) uruchomień.
 - Narzędzia:** Ikony akcji (uruchom, edytuj, usuń, itp.).
- Przycisk + Nowy Przepływ Pracy:** Otwiera okno modalne do konfiguracji nowego przepływu pracy od zera.
- Przycisk Użyj szablonu:** To najszybszy sposób na stworzenie zaawansowanej automatyzacji. Zamiast konfigurować każdy krok ręcznie, możesz skorzystać z biblioteki gotowych do użycia przepływów pracy, przygotowanych dla typowych zastosowań biznesowych.

Wybierz Szablon Przepływu Pracy

Finanse i Księgowość

Automatyczne Przetwarzanie Faktur Kosztowych

Monitoruje dedykowaną skrzynkę e-mail, importuje załączone faktury i automatycznie wyodrębnia z nich kluczowe dane (NIP, kwoty, daty) za pomocą Agenta AI.

Dział Prawny i Compliance

Wstępna Analiza Umów NDA

Automatycznie skanuje wskazany folder w poszukiwaniu nowych umów o poufności (NDA), a następnie uruchamia Agenta Prawnego, który analizuje dokument pod kątem kluczowych klauzul (kary umowne, okres obowiązywania itp.).

HR / Kadry

Automatyczny Screening CV

Monitoruje skrzynkę rekrutacyjną, pobiera załączone CV, a następnie prosi Agenta HR o wstępną weryfikację kompetencji kandydata pod kątem konkretnego stanowiska.

- **Cel i korzyści:**

- **Przyspieszenie pracy:** Wybór szablonu natychmiast wypełnia cały formularz wyzwalaczem i sekwencją akcji.
- **Inspiracja i dobre praktyki:** Szablony są praktycznymi przykładami pokazującymi, jak efektywnie łączyć różne komponenty systemu (np. import z e-maila, analiza przez Agenta AI i wysyłka powiadomienia).
- **Minimalizacja błędów:** Gotowe konfiguracje są przetestowane i zoptymalizowane, co zmniejsza ryzyko pomyłki.

- **Jak to działa?**

1. Po kliknięciu Użyj szablonu pojawia się lista szablonów pogrupowanych w kategorie (np. *Finanse i Księgowość, Dział Prawny, HR*).
2. Każdy szablon ma nazwę i opis (np. *Automatyczne Przetwarzanie Faktur Kosztowych*).
3. Wybranie szablonu otwiera standardowe okno konfiguracji, ale wszystkie pola są już wypełnione. Twoim zadaniem jest jedynie dostosowanie wartości (np. zmiana adresu e-mail czy treści polecenia dla Agenta) i zapisanie gotowego przepływu.

4. **Akcje (ikony przy każdym zadaniu):**



- **Uruchom teraz:** Natychmiastowe, jednorazowe uruchomienie zadania, niezależnie od harmonogramu.
- **Edytuj:** Otwiera okno z pełną konfiguracją przepływu.
- **Usuń:** Trwale usuwa zadanie automatyzacji.
- **Ponów Błędy:** Przenosi pliki z podkatalogu failed z powrotem do folderu źródłowego w celu ponownej próby przetworzenia.
- **Dziennik Zdarzeń:** Otwiera historię uruchomień zadania, pokazując status, znaną liczbę plików i czas trwania każdego cyklu.

Dziennik Zdarzeń (Workflow ID: 1)



Data	Status	Podsumowanie	Czas
1.10.2025, 12:14:03	SUCCESS	Zakończono skanowanie. Zakolejkowano 0 z 0 znalezionych plików.	0.01s
1.10.2025, 12:13:03	SUCCESS	Zakończono skanowanie. Zakolejkowano 0 z 0 znalezionych plików.	0.01s
1.10.2025, 12:12:03	SUCCESS	Zakończono skanowanie. Zakolejkowano 0 z 0 znalezionych plików.	0.01s
1.10.2025, 12:11:03	SUCCESS	Zakończono skanowanie. Zakolejkowano 0 z 0 znalezionych plików.	0.00s

Tworzenie i Konfiguracja Przepływu Pracy

Edytuj Przepływ Pracy

Podstawowe informacje

Nazwa

Upload plików

Organizacja

EnterSoft

Opis

☒ Aktywny przepływ pracy

Wyzwalacz (Kiedy uruchomić?)

Harmonogram

Nowy E-mail

Uruchom co (minuty)

1

Akcje (Co zrobić?)

Krok 1: Importuj Pliki

Wstaw domyślne

Szafka katalogu

EnterSoft/Dolimportu

Ilość plików na cykl

10

Timeout dla pliku (sekundy)

300

Polityka współbieżności

Zakoleguj

Wzorzec nazwy pliku

Domyślne tagi (JSON)

[]

Wstaw przykład

Po sukcesie:

Przenieś

Po błędzie:

Przenieś

☒ Skanuj rekursywnie

☒ Nadpisz duplikaty

☒ Głęboka analiza OCR

Email (Sukces)

jacek@entersoft.pl

Email (Błąd)

biuro@entersoft.pl

Krok 2: Uruchom Agenta

Wstaw domyślne

Nazwa Agenta

Agent ds. Streszczeń i Zadań

Polecenie dla Agenta

Na podstawie dostarczonego kontekstu, streść załączony dokument i wypunktuj jego najważniejsze zagadnienia w formacie Markdown.

Krok 3: Wyślij Email

Wstaw domyślne

Adres e-mail odbiorcy

jacek@entersoft.pl

Szablon tematu

[AsWiseAI] Wynik analizy dla pliku: {file_name}

Szablon treści

Dzień dobry,

Agent AI zakończył analizę dokumentu. Poniżej znajdują się szczegóły i wynik.

Szczegóły zlecenia:

- Plik: {file_name}

- Agent: {agent_name}

- Workflow: {workflow_name}

Wynik analizy:

{agent_result}

Wiadomość wygenerowana automatycznie przez system AsWiseAI. {timestamp}

Krok 4: Wyślij Powiadomienie

Wstaw domyślne

Szablon powiadomienia

☒ Analiza pliku {file_name} przez agenta {agent_name} została zakończona.

Dostępne zmienne: {file_name}, {agent_name}

Krok 5: Zapisz Wynik do Pliku

Wstaw domyślne

Katalog docelowy

EnterSoft/WynikiAnalizy

Szablon nazwy pliku

Raport_{file_name}_{timestamp}.md

Szablon treści pliku

Raport Analizy - AsWiseAI

Przebieg pracy: {workflow_name}

Przeanalizowany plik: {file_name}

Agent wykonujący: {agent_name}

Data analizy: {timestamp}

AsWiseAI, AsWiseAI, AsWiseAI, RT

Importuj Pliki

Uruchom Agenta

Wyślij Email

Wyślij Powiadomienie

Zapisz do Pliku

Proces konfiguracji składa się z trzech głównych części:

1. Konfiguracja Podstawowa



EnterSoft

Firma Informatyczna EnterSoft

e-mail: biuro@entersoft.pl, entersoft.pl, kom.: 696 481 842

Str.: 46/99

Nowy Przepływ Pracy

Podstawowe informacje

Nazwa

Organizacja

EnterSoft

Opis

☒

 Aktywny przepływ pracy

- **Nazwa i Opis:** Nadaj przepływowi czytelną nazwę i opcjonalnie opisz jego przeznaczenie.
- **Organizacja:** Przypisz przepływ do konkretnej organizacji.
- **Zadanie aktywne:** Główny włącznik/wyłącznik dla danego zadania.

2. Konfiguracja Wyzwalacza (Trigger)

Wyzwalacz to zdarzenie, które inicjuje cały przepływ pracy.

- **Uruchomienie Czasowe (Harmonogram)**

Wyzwalacz (Kiedy uruchomić?)

Harmonogram

Nowy E-mail

Uruchom co (minuty)

10

Przepływ będzie uruchamiany cyklicznie w określonym interwale (np. co 15 minut).

- **Wyzwalacz E-mail**

Wyzwalacz (Kiedy uruchomić?)

Harmonogram

Nowy E-mail

Serwer IMAP

Port IMAP

993

Użytkownik

Hasło

..... (wpisz, aby zmienić)

☐ Przetwarzaj załączniki

☐ Przetwarzaj treść e-maila

System monitoruje wskazaną skrzynkę pocztową IMAP. Przepływ jest uruchamiany dla każdego nowego, maila bez etykiety „AsWiseAI_Processed”.

Możesz skonfigurować, czy system ma przetwarzać **załączniki**, **treść wiadomości** (jako plik .txt), czy oba.

Wyzwalacz E-mail to potężny mechanizm, który przekształca Twoją skrzynkę pocztową w zautomatyzowany punkt wejścia dla dokumentów i informacji do systemu AsWiseAI. Zamiast ręcznie pobierać pliki i przysyłać je do aplikacji, możesz skonfigurować przepływ pracy tak, aby sam monitorował dedykowany adres e-mail i automatycznie rozpoczynał przetwarzanie, gdy tylko pojawi się nowa wiadomość.

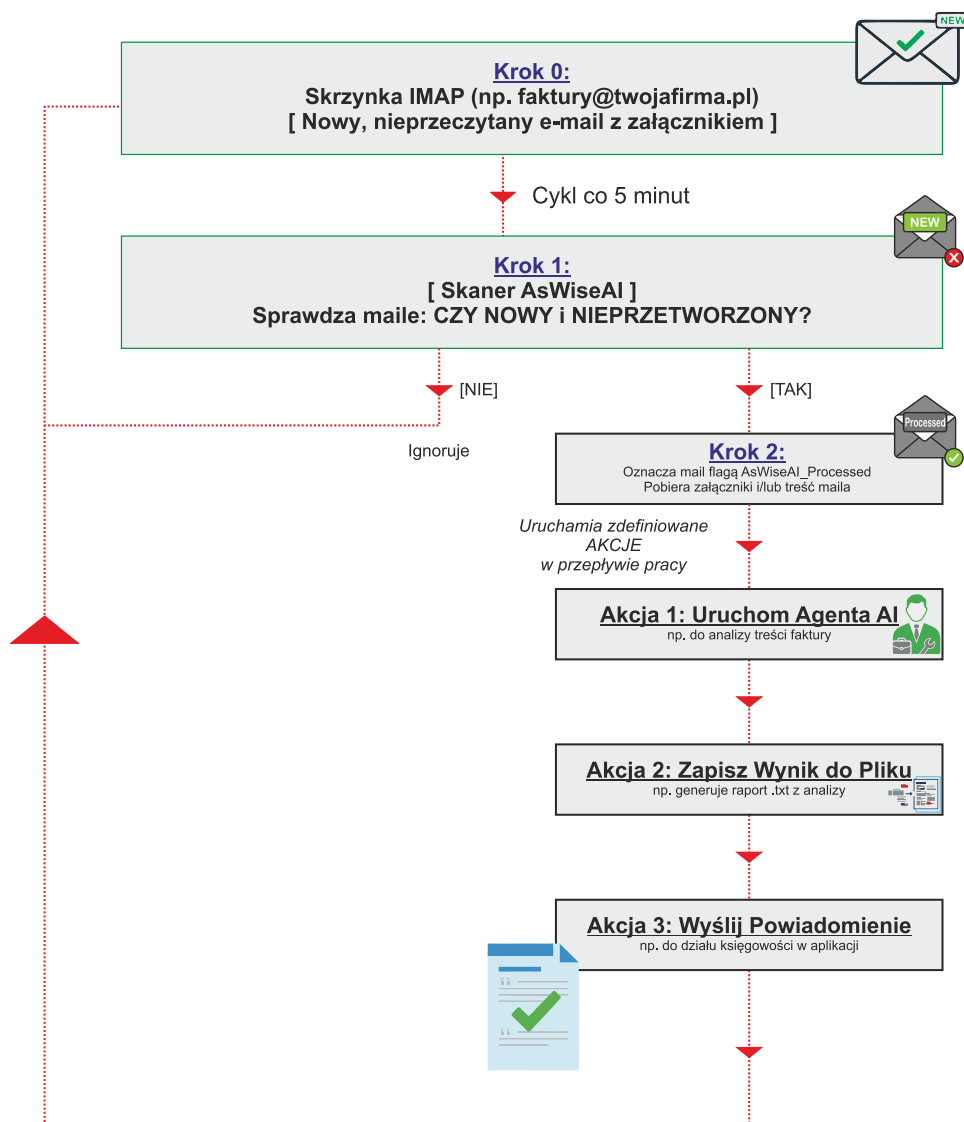
Jak to Działa? Mechanizm Krok po Kroku

1. **Połączenie ze Skrzynką (IMAP):** W tle, w regularnych odstępach czasu (zwykle co 5 minut), system AsWiseAI łączy się ze wskazaną przez Ciebie skrzynką pocztową przy użyciu bezpiecznego protokołu IMAP.
2. **Wyszukiwanie Nowych Wiadomości:** System nie pobiera wszystkich maili. Aby zapewnić wydajność i unikać wielokrotnego przetwarzania, wyszukuje on wyłącznie wiadomości spełniające **dwa warunki jednocześnie**:
 - Są **nieprzeczytane** (flaga UNSEEN).
 - **NIE posiadają** specjalnej, niestandardowej flagi AsWiseAI_Processed.
3. **Oznaczanie Przetworzonych Wiadomości:** Gdy tylko system zidentyfikuje wiadomość do przetworzenia, natychmiast oznacza ją na serwerze dwiema flagami:
 - \Seen (jako przeczytaną).
 - AsWiseAI_Processed (jako przetworzoną przez system). Dzięki temu ten sam e-mail **nigdy nie zostanie przetworzony ponownie**, nawet jeśli ktoś ręcznie oznaczy go z powrotem jako nieprzeczytany.
4. **Ekstrakcja Danych:** W zależności od konfiguracji, system pobiera z wiadomości:
 - **Załączniki:** Wszystkie załączone pliki (PDF, obrazy, TXT) są pobierane i traktowane jako osobne dokumenty do przetworzenia.
 - **Treść E-maila:** Treść samej wiadomości jest zapisywana jako tymczasowy plik `trecsc_maila.txt` i również traktowana jako dokument do przetworzenia.
5. **Uruchomienie Sekwencji Akcji:** Pobrane w ten sposób pliki (załączniki i/lub treść) są przekazywane do pierwszej **Akcji** zdefiniowanej w Twoim przepływie pracy, rozpoczynając całą sekwencję automatyzacji.



Uproszczony schemat działania przepływu pracy Workflow e-mail

Przepływy Pracy to Twój osobisty automat, który monitoruje źródła danych, zleca analizy Agentom AI i dostarcza wyniki tam, gdzie ich potrzebujesz (np. na e-mail lub jako gotowy raport).



❖ Możliwe Zastosowania i Przykłady

Wyzwalacz E-mail jest niezwykle wszechstronny. Oto kilka typowych scenariuszy biznesowych, w których sprawdza się idealnie:

1. Automatyzacja Faktur Kosztowych

- **Scenariusz:** Dział księgowości otrzymuje faktury od dostawców w formacie PDF na dedykowany adres e-mail, np. faktury@twojafirma.pl.
- **Konfiguracja Wyzwalacza:**
 - **Serwer IMAP:** imap.twojafirma.pl

- **Użytkownik:** faktury@twojafirma.pl
- **Przetwarzaj załączniki:** Włączone.
- **Przetwarzaj treść:** Wyłączone.
- **Przykładowe Akcje:**
 1. **Uruchom Agent AI:** Agent analityczny otrzymuje PDF-a i dostaje polecenie: *"Wyodrębnij NIP sprzedawcy, numer faktury, datę wystawienia i kwotę brutto. Zwróć wynik w formacie JSON."*
 2. **Wyślij E-mail:** Wynik analizy ({agent_result}) jest wysyłany do osoby odpowiedzialnej za księgowania w celu weryfikacji.

2. Automatyczny Screening CV

- **Scenariusz:** Dział HR prowadzi rekrutację na stanowisko "Senior Python Developer" i otrzymuje CV kandydatów na adres rekrutacja-python@twojafirma.pl.
- **Konfiguracja Wyzwalacza:**
 - **Serwer IMAP:** imap.twojafirma.pl
 - **Użytkownik:** rekrutacja-python@twojafirma.pl
 - **Przetwarzaj załączniki:** Włączone.
- **Przykładowe Akcje:**
 1. **Uruchom Agent AI:** Agent HR otrzymuje CV i polecenie: *"Przeanalizuj to CV pod kątem oferty pracy na stanowisko 'Senior Python Developer'. Wylistuj technologie wymienione w dokumencie i oszacuj lata doświadczenia komercyjnego."*
 2. **Zapisz Wynik do Pliku:** Podsumowanie od Agent AI jest zapisywane do pliku [CV]_{file_name}.txt w dedykowanym folderze na serwerze.
 3. **Wyślij Powiadomienie:** Rekruter otrzymuje w aplikacji AsWiseAI powiadomienie: *"Otrzymano i przeanalizowano nowe CV: {file_name}."*

3. Kategoryzacja Zgłoszeń Supportowych

- **Scenariusz:** Klienci wysyłają zgłoszenia techniczne na adres support@twojafirma.pl. Chcemy automatycznie przypisać im priorytet i kategorię.
- **Konfiguracja Wyzwalacza:**
 - **Serwer IMAP:** imap.twojafirma.pl
 - **Użytkownik:** support@twojafirma.pl
 - **Przetwarzaj załączniki:** Wyłączone.
 - **Przetwarzaj treść:** Włączone.
- **Przykładowe Akcje:**

1. **Uruchom Agenta AI:** Agent ds. Zgodności otrzymuje treść maila i polecenie: *"Na podstawie treści zgłoszenia przypisz jedną z kategorii: 'Problem z Logowaniem', 'Błąd Aplikacji', 'Zapytanie Ogólne' oraz jeden z priorytetów: 'Niski', 'Normalny', 'Wysoki'. Odpowiedz w formacie JSON."*
2. **Wyślij E-mail:** Wynik kategoryzacji jest przesyłany do systemu ticketowego lub bezpośrednio do lidera zespołu wsparcia.

3. Konfiguracja Akcji

Akcje to sekwencyjne kroki, które system wykonuje po uruchomieniu przepływu. Możesz dodać wiele akcji, które będą wykonywane jedna po drugiej.

+ Importuj Pliki + Uruchom Agenta + Wyślij Email + Wyślij Powiadomienie + Zapisz do Pliku

Dostępne typy akcji:

- **Importuj Pliki:**

Akcje (Co zrobić?) ⓘ

Krok 1: Importuj Pliki ▼

Wstaw domyślne

Ścieżka katalogu ⓘ
EnterSoft/Dolimportu

Ilość plików na cykl
10

Timeout dla pliku (sekundy)
300

Polityka współbieżności
Pomiń ▼

Wzorzec nazwy pliku ⓘ
*.pdf

Domyślne tagi (JSON) ⓘ
[]

Wstaw przykład

Po sukcesie:
Przenieś ▼

Po błędzie:
Przenieś ▼

☒ Skanuj rekursywnie

☒ Nadpisz duplikaty

☒ Głęboka analiza OCR

Email (Sukces)

Email (Błąd)

- **Ścieżka Katalogu:** Względna ścieżka do katalogu, który ma być monitorowany (wewnątrz /app/data/auto_ingest/).
- **Wzorzec nazwy pliku:** (Opcjonalne) Pozwala przetwarzać tylko pliki pasujące do wzorca, np. FAKTURA_*.pdf.

- **Ilość plików na cykl:** Maksymalna liczba plików przetwarzana w jednym cyklu.
 - **Domyślne Tagi (JSON):** Lista tagów automatycznie przypisywanych do każdego pliku, np. ["faktura", "kosztowa"].
 - **Po sukcesie / Po błędzie:** Określa, co ma się stać z plikiem źródłowym po przetworzeniu (przenies do processed/failed, usuń lub zostaw).
 - **Opcje dodatkowe:** Skanowanie rekursywne, nadpisywanie duplikatów czy włączenie głębokiej analizy OCR.
- **Uruchom Agent AI:**

Krok 2: Uruchom Agent AI Wstaw domyślne

Nazwa Agent AI

Agent ds. Streszczeń i Zadań

Polecenie dla Agent AI

Jesteś analitykiem projektu. Przeanalizuj ten raport statusowy i zwróć odpowiedź w formacie HTML. Twoja analiza musi zawierać następujące sekcje:

```

### Podsumowanie menedżerskie
(1-2 zdania opisujące ogólny stan projektu)

### Zidentyfikowane Ryzyka
(Wypunktowana lista wszystkich ryzyk wspomnianych w raporcie, wraz z ich priorytetem, jeśli jest podany)

### Zadania dla Zespołu Backend
(Wypunktowana lista zadań przypisanych bezpośrednio do zespołu 'backend')

```

- Pozwala na automatyczne uruchomienie jednego z Twoich Agentów AI (zdefiniowanych wcześniej w Macierzy Agentów AI).
 - Agent przetwarza treść pliku z poprzedniego kroku na podstawie zdefiniowanego przez Ciebie polecenia (promptu).
 - *Przykład:* Możesz polecić Agentowi Analitykowi Danych: "Z tego dokumentu wyodrębnij NIP sprzedawcy, numer faktury i kwotę brutto. Zwróć wynik w formacie JSON."
- **Wyślij E-mail:**

Krok 3: Wyślij Email ▼ Wstaw domyślne 🗑️

Adres e-mail odbiorcy

biuro@entersoft.pl

Szablon tematu ⓘ

[AsWiseAI] Wynik analizy dla pliku: {file_name}

Szablon treści ⓘ

Dzień dobry,

Agent AI zakończył analizę dokumentu.
Poniżej znajdują się szczegóły i wynik.

Szczegóły zlecenia:

- Plik: {file_name}
- Agent: {agent_name}
- Workflow: {workflow_name}

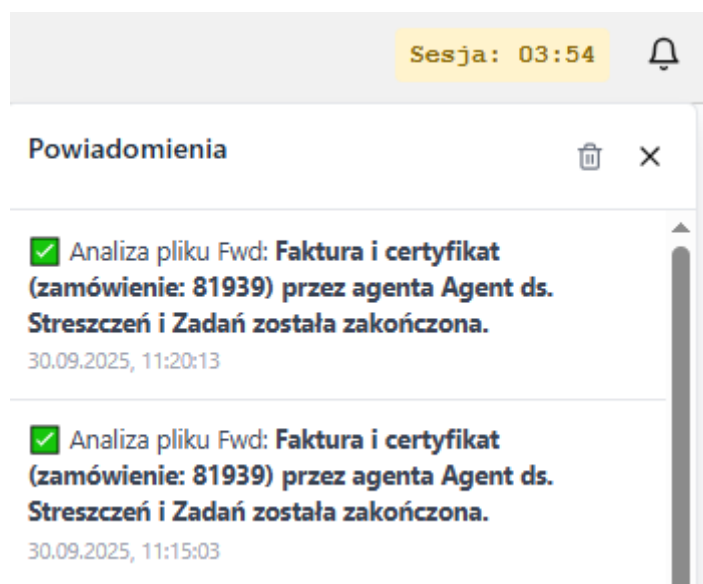
Wynik analizy:

{agent_result}

Wiadomość wygenerowana automatycznie przez system AsWiseAI.

- Wysła wiadomość e-mail z podsumowaniem. Możesz dynamicznie wstawiać do treści i tematu dane z poprzednich kroków, używając placeholderów, np. {file_name}, {agent_result}.

- **Wyślij Powiadomienie:**



- Tworzy powiadomienie wewnątrz aplikacji AsWiseAI dla wskazanych użytkowników.

- **Zapisz Wynik do Pliku:**

Krok 7: Zapisz Wynik do Pliku

Wstaw domyślne

Katalog docelowy

EnterSoft/WynikiAnalizy

Ścieżka względna do /app/data/auto_ingest.

Szablon nazwy pliku

Raport_{file_name}_{timestamp}.md

Szablon treści pliku

```

# Raport Analizy - AswiseAI
---
- **Przepływ pracy:** {workflow_name}
- **Przeanalizowany plik:** {file_name}
- **Agent wykonujący:** {agent_name}
- **Data analizy:** {timestamp}
---

## Wynik Analizy Agent AI

{agent_result}

```

- Zapisuje wynik działania Agent AI do nowego pliku tekstowego we wskazanym katalogu. Nazwę pliku i jego treść można również dynamicznie generować za pomocą placeholderów.

Dziennik Zdarzeń i Monitoring

Każde uruchomienie przepływu pracy jest odnotowywane w **Dzienniku Zdarzeń**. Znajdziesz tam szczegółowe informacje o:

- **Statusie:** SUCCESS (sukces), FAILURE (błąd), RUNNING (w trakcie).
- **Podsumowaniu:** Liczba znalezionych, zakolejkowanych i pominiętych plików.
- **Czasie trwania:** Jak długo trwał dany cykl.

Jeśli jakiś plik nie zostanie przetworzony poprawnie, trafi do podkatalogu failed. Możesz przeanalizować przyczynę błędu w logach systemowych, a następnie użyć przycisku **Ponów Błąd**, aby spróbować przetworzyć go ponownie.

Archiwizacja AI

 **Agent Archiwizujący**

Zarządzaj automatycznym, cyklicznym zadaniem, które analizuje stare dokumenty i oznacza je do archiwizacji, jeśli uzna ich treść za nieaktualną.

Harmonogram Archiwizacji

Ostatnie uruchomienie: **Nigdy**

Minuta

0

Godzina

2

Dzień miesiąca

*

Dzień tygodnia

*

 **Zapisz**

Użyj "*" dla "każdy" lub składni CRON, np. "*/15" co 15 minut, "1,15" dla 1 i 15 dnia miesiąca.

Zakładka **Archiwizacja AI** to narzędzie służące do inteligentnego zarządzania cyklem życia dokumentów w Twojej bazie wiedzy. Zamiast ręcznie przeglądać setki lub tysiące starych plików, możesz zlecić to zadanie wyspecjalizowanemu agentowi AI, który pomoże zidentyfikować nieaktualne już treści.

Cel i korzyści

Głównym celem modułu Archiwizacji AI jest utrzymanie bazy wiedzy w aktualnym stanie. Regularne oznaczanie przestarzałych dokumentów przynosi trzy kluczowe korzyści:

1. **Zwiększa trafność odpowiedzi AI:** Zapobiega sytuacji, w której modele językowe opierają swoje odpowiedzi na nieaktualnych danych, starych procedurach czy nieobowiązujących już regulaminach.
2. **Ułatwia porządkowanie danych:** Automatycznie kategoryzuje dokumenty, które są kandydatami do przeniesienia do archiwum długoterminowego lub do usunięcia.
3. **Oszczędza czas administratorów:** Zamiast manualnie przeglądać dokumenty, administratorzy otrzymują listę plików wstępnie zakwalifikowanych przez AI do archiwizacji.

Jak to działa? Proces Krok po Kroku

Proces archiwizacji jest w pełni zautomatyzowany i opiera się na działaniu dedykowanego **Agent** **Archiwizującego**.

1. **Cykliczne Skanowanie:** Zgodnie z ustawionym przez Ciebie harmonogramem, system automatycznie przeszukuje bazę wiedzy w poszukiwaniu dokumentów, które są starsze niż zdefiniowany próg (domyślnie 2 lata) i nie zostały jeszcze oznaczone do archiwizacji.
2. **Analiza przez Agent** **AI:** Każdy zidentyfikowany jako "stary" dokument jest przekazywany do **Agent** **Archiwizującego**. Agent ten analizuje treść i metadane dokumentu, starając się zrozumieć jego kontekst i ocenić, czy informacje w nim zawarte są nadal aktualne. Agent poszukuje oznak nieaktualności, takich jak:
 - Daty wydarzeń, które już minęły.
 - Zakończone promocje lub oferty.

- Stare wersje regulaminów, procedur czy polityk.
 - Odniesienia do nieistniejących już projektów lub produktów.
3. **Inteligentne Oznaczanie:** Jeśli Agent Archiwizujący uzna, że dokument jest nieaktualny, nie usuwa go. Zamiast tego, dodaje do wszystkich jego fragmentów w bazie wiedzy specjalny tag: **[do-archiwizacji]**.

Co dzieje się po oznaczeniu?

Tag [do-archiwizacji] sam w sobie nie wykonuje żadnej destrukcyjnej akcji. Jest to inteligentna etykieta dla administratora. Po zakończeniu cyklu archiwizacji możesz:

1. Przejść do panelu **Baza Wiedzy**.
2. Użyć filtra tagów, aby wyświetlić wszystkie dokumenty oznaczone jako [do-archiwizacji].
3. Przejrzeć listę i podjąć ostateczną decyzję o ich usunięciu lub przeniesieniu do zewnętrznego archiwum.

Konfiguracja Panelu

Panel Archiwizacji AI pozwala na pełną kontrolę nad procesem:

- **Główny Włącznik:** Umożliwia włączenie lub wyłączenie całego mechanizmu archiwizacji.
- **Harmonogram (CRON):** Zestaw pól (Minuta, Godzina, Dzień miesiąca itd.) pozwala precyzyjnie zdefiniować, kiedy zadanie ma być uruchamiane. Domyślnie jest to **każdego pierwszego dnia miesiąca o godzinie 2:00 w nocy**, aby nie obciążać systemu w godzinach pracy.
- **Zapisz Zmiany:** Przycisk do zatwierdzenia nowej konfiguracji harmonogramu i statusu.
- **Ostatnie Uruchomienie:** Wyświetla datę i godzinę ostatniego wykonania zadania, co pozwala monitorować jego regularne działanie.

Integracje (LDAP / Active Directory)

Panel **Integracje** (dostępny w Ustawieniach Systemowych) umożliwia skonfigurowanie centralnego uwierzytelniania użytkowników za pomocą protokołu LDAP, najczęściej w celu integracji z Microsoft Active Directory.

Ustawienia Systemowe

Automatyzacja Integracje Reguły AI Kopie Zapasowe Reset Danych

Zmiany w tej sekcji wpływają na sposób logowania użytkowników. Konfiguruj ostrożnie i zawsze testuj połączenie przed zapisaniem.

☒ Integracja z Active Directory

URI Serwera LDAPs

ldaps://dc.twojadamena.pl:636

Bind DN (konto serwisowe)

CN=svc_aswiseai,OU=Services,DC=...

Hasło Konta Serwisowego

Wprowadź hasło

Search Base DN (ścieżka do użytkowników)

OU=Users,DC=twojadamena,DC=pl

Mapowanie Grup AD na Role

NowaGrupaAD_0	→	User	
NowaGrupaAD_1	→	Admin	

[Dodaj mapowanie](#)

[Testuj Połączenie](#) [Zapisz Zmiany](#)

(Rys.: Panel konfiguracji integracji z LDAP / Active Directory)

Cel integracji: Po włączeniu i poprawnym skonfigurowaniu integracji, użytkownicy mogą logować się do AsWiseAI za pomocą swoich standardowych danych domenowych (nazwy użytkownika i hasła), bez potrzeby tworzenia dla nich osobnych kont w aplikacji.

Proces działania

1. Użytkownik próbuje zalogować się do AsWiseAI swoimi danymi domenowymi.
2. AsWiseAI łączy się z serwerem Active Directory i weryfikuje poprawność danych.
3. Jeśli uwierzytelnienie w AD się powiedzie, AsWiseAI automatycznie tworzy (lub aktualizuje) konto dla tego użytkownika w swojej lokalnej bazie danych, przypisując mu rolę na podstawie mapowania grup.
4. Użytkownik zostaje zalogowany do systemu.

Tworzenie konta użytkownika w AsWiseAI

Gdy użytkownik po raz pierwszy loguje się pomyślnie za pomocą swoich danych z Active Directory, AsWiseAI automatycznie tworzy dla niego konto w lokalnej bazie. Ten wpis przechowuje informacje o nazwie użytkownika, jego roli i przypisaniu do organizacji, ale celowo nie przechowuje hasła. Każde kolejne logowanie jest weryfikowane bezpośrednio z serwerem AD, co zapewnia maksymalne bezpieczeństwo i centralne zarządzanie dostępem.

Pola konfiguracyjne

- **URI Serwera LDAPS:** Pełny adres serwera Active Directory, np. ldaps://dc.twojadenomena.pl:636. Zalecane jest użycie bezpiecznego połączenia ldaps.
- **Bind DN (konto serwisowe):** Pełna nazwa konta technicznego, którego AsWiseAI użyje do wyszukiwania użytkowników w AD.
- **Hasło Konta Serwisowego:** Hasło dla powyższego konta. Jest ono przechowywane w bazie w formie zaszyfrowanej.
- **Search Base DN:** Ścieżka w strukturze AD, od której system ma zacząć wyszukiwanie użytkowników (np. OU=Uzytkownicy,DC=twojadenomena,DC=pl).
- **Mapowanie Grup AD na Role:** Pozwala na automatyczne przypisanie roli w AsWiseAI (user lub admin) na podstawie przynależności użytkownika do grupy w Active Directory. W lewym polu należy wpisać dokładną nazwę grupy z AD, a w prawym wybrać rolę.
- **Testuj Połączenie:** Przycisk pozwalający zweryfikować poprawność wprowadzonych danych bez ich zapisywania.

Jak działa logowanie po włączeniu integracji?

Włączenie integracji z Active Directory nie blokuje możliwości logowania dla kont lokalnych. System działa w trybie priorytetowym:

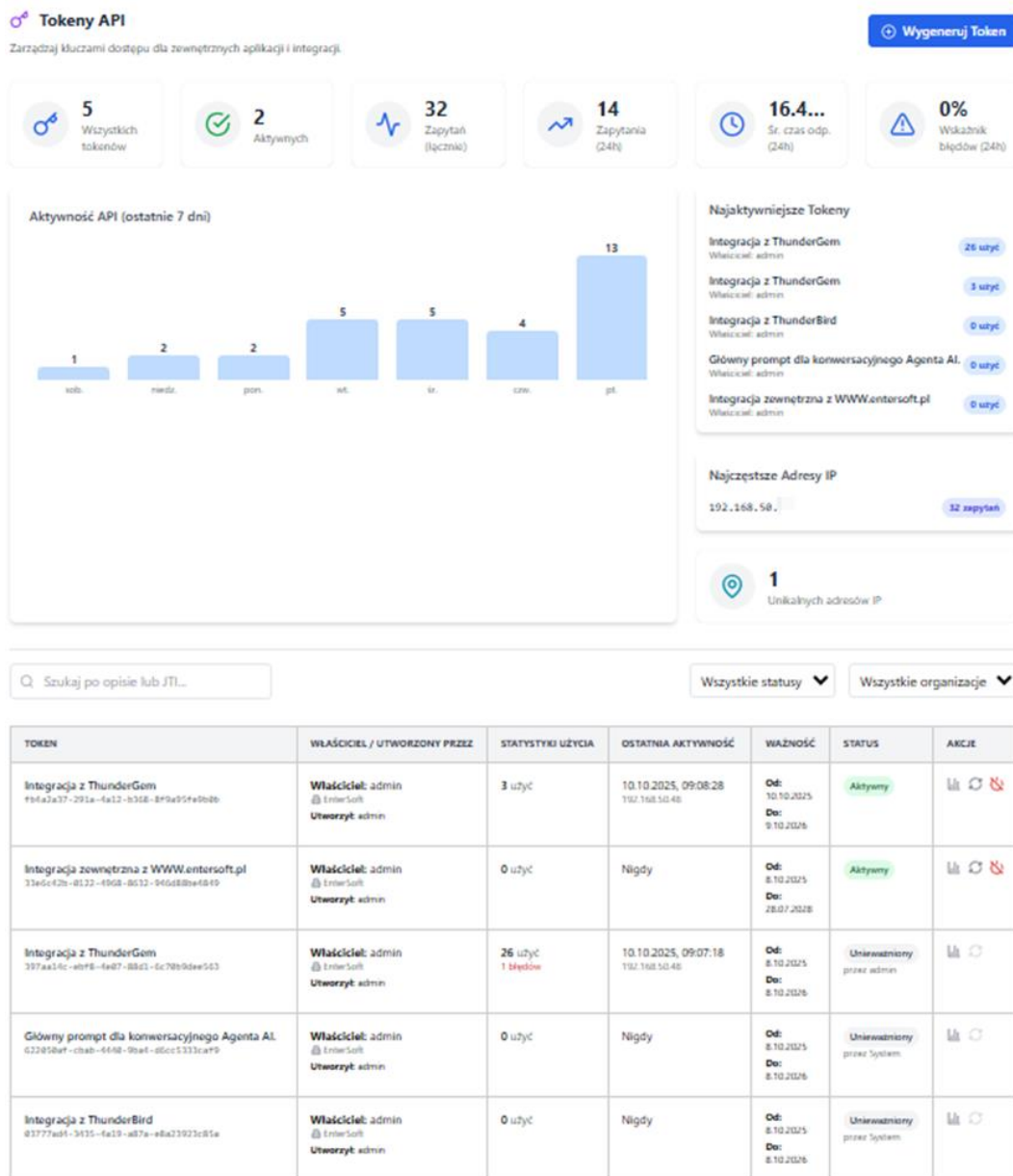
- **Logowanie domenowe (priorytet):** Gdy użytkownik wpisuje swoje dane, AsWiseAI najpierw próbuje go uwierzytelnić na serwerze Active Directory. Jeśli się to powiedzie, użytkownik uzyskuje dostęp.
- **Logowanie lokalne (fallback):** Jeśli próba logowania przez AD nie powiedzie się, system automatycznie podejmuje drugą próbę, sprawdzając, czy dane pasują do konta zapisanego w lokalnej bazie danych AsWiseAI.

W praktyce oznacza to, że:

- Użytkownicy istniejący w Active Directory mogą logować się swoimi standardowymi danymi firmowymi.
- Konta, które istnieją **tylko** w AsWiseAI (jak domyślne konto admin lub inne konta techniczne), **nadal mogą się logować** za pomocą swoich lokalnych haseł.

Wyłączenie przełącznika w panelu powoduje, że system całkowicie pomija próbę logowania przez AD i od razu weryfikuje użytkowników tylko w lokalnej bazie.

Tokeny API (Dostęp dla Aplikacji Zewnętrznych)



Moduł **Tokeny API** to zaawansowana funkcja przeznaczona dla **Super Administratorów**, umożliwiająca tworzenie bezpiecznych kluczy dostępu (tokenów) dla zewnętrznych aplikacji i integracji. Dzięki nim możliwe jest połączenie z AsWiseAI innych systemów, takich jak firmowe CRM, aplikacje analityczne, klienci poczty (np. ThunderBird z dodatkiem [ThunderGem](#)) czy niestandardowe skrypty, w celu automatycznego zadawania pytań i pobierania odpowiedzi bez konieczności logowania przez interfejs użytkownika.

Mechanizm ten bazuje na **stanowych, długowiecznych tokenach JWT**, których ważność i uprawnienia są weryfikowane przy każdym zapytaniu bezpośrednio w bazie danych aplikacji, co zapewnia najwyższy poziom bezpieczeństwa i kontroli.

Zasady Działania i Schemat Przepływu Danych



Każdy wygenerowany token jest trwale powiązany z konkretnym użytkownikiem w ramach jego organizacji. Oznacza to, że aplikacja korzystająca z tokena działa dokładnie z takimi samymi uprawnieniami, jakie posiada ten użytkownik. Gwarantuje to, że zewnętrzna integracja nigdy nie uzyska dostępu do danych, do których użytkownik-właściciel tokena nie miałby dostępu.

Kluczowe cechy mechanizmu tokenów

- **Stanowość:** Każdy token posiada swój unikalny identyfikator zapisany w bazie danych, co pozwala na jego śledzenie i unieważnienie w dowolnym momencie.

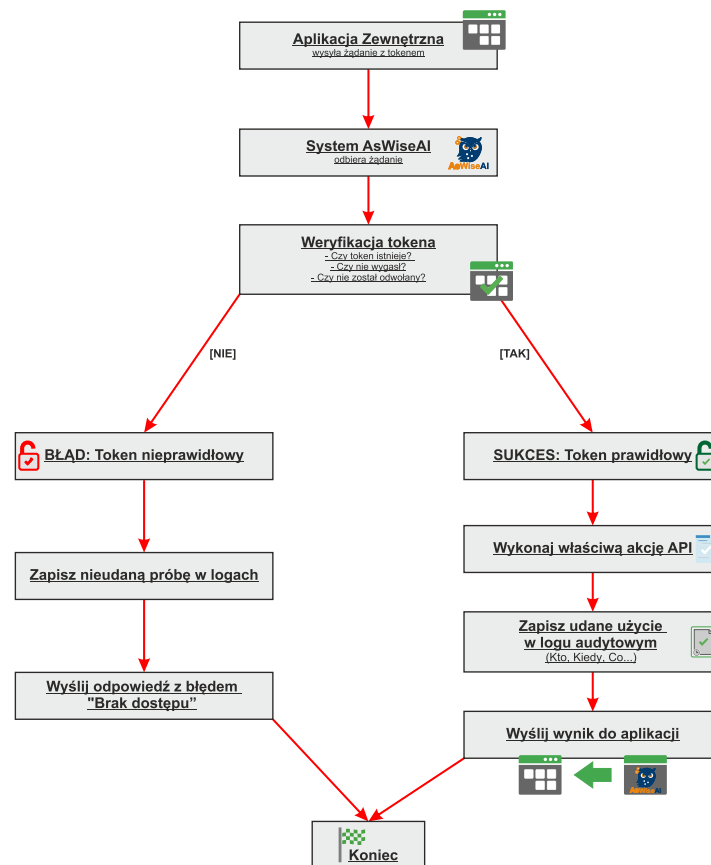
- **Bezpieczeństwo:** Nawet jeśli token zostanie przechwycony, Super Administrator może go natychmiast unieważnić, odcinając dostęp do systemu. System przy każdym użyciu sprawdza, czy token nie został odwołany lub czy nie wygasł.
- **Uprawnienia:** Podczas tworzenia tokena można mu nadać tylko te uprawnienia, które posiada wybrany użytkownik. Zapewnia to zgodność z zasadą minimalnych uprawnień - token otrzymuje dostęp tylko do tych operacji, które są niezbędne do jego działania.
- **Pełna Ścieżka Audytowa:** System przechowuje szczegółowe informacje o każdym tokenie oraz loguje każde jego użycie, co umożliwia pełną weryfikację jego historii i aktywności.

Schemat Uwierzytelniania i Audytu Tokena



Schemat Uwierzytelniania i Audytu Tokena

Poniższy schemat ilustruje, co dzieje się za każdym razem, gdy zewnętrzna aplikacja wysła zapytanie do AsWiseAI przy użyciu tokena API



Powyższy schemat ilustruje, co dzieje się za każdym razem, gdy zewnętrzna aplikacja wysła zapytanie do AsWiseAI przy użyciu tokena API.

1. **Żądanie z Tokenem:** Zewnętrzna aplikacja wysyła zapytanie do API AsWiseAI, dołączając token w nagłówku.
2. **Zapytanie jest przechwytywane** przez warstwę pośredniczącą, która rozpoczyna pomiar czasu wykonania.
3. **Weryfikacja w Bazie Danych:** System dekoduje token, odczytuje jego unikalny identyfikator i sprawdza w bazie danych, czy token jest aktywny (nie wygasł i nie został unieważniony).
4. **Aktualizacja Statystyk:**



- **Sukces:** Jeśli token jest ważny, system zwiększa licznik użyc i aktualizuje datę ostatniego użycia oraz adres IP.
 - **Błąd:** Jeśli token jest nieważny, zwiększany jest licznik nieudanych prób.
5. **Logowanie Użycia:** Po wykonaniu żądania, warstwa pośrednia zatrzymuje stoper i zapisuje szczegółowy wpis w dzienniku zdarzeń, zawierający m.in. wywołany endpoint, kod odpowiedzi i czas trwania operacji.
 6. **Odpowiedź:** Klient API otrzymuje odpowiedź.

Panel Zarządzania Tokenami

Panel dostępny w Ustawienia Systemowe > Tokeny.

Wszystkie statusy

Wszystkie organizacje

TOKEN	WŁAŚCICIEL / UTWORZONY PRZEZ	STATYSTYKI UŻYCIA	OSTATNIA AKTYWNOŚĆ	WAŻNOŚĆ	STATUS	AKCJE
Integracja z ThunderGem fb4e6c42b-291a-4a12-b368-8f9d95fe9b0b	Właściciel: admin EnterSoft Utworzył: admin	3 użyc	10.10.2025, 09:08:28 192.168.50.48	Od: 10.10.2025 Do: 9.10.2026	Aktywny	
Integracja zewnętrzna z WWW.entersoft.pl 33e6c42b-0122-4968-8632-946d88be4849	Właściciel: admin EnterSoft Utworzył: admin	0 użyc	Nigdy	Od: 8.10.2025 Do: 28.07.2028	Aktywny	
Integracja z ThunderGem 397aa14c-ebf8-4e07-88d1-6c70b9dee563	Właściciel: admin EnterSoft Utworzył: admin	26 użyc 1 błądów	10.10.2025, 09:07:18 192.168.50.48	Od: 8.10.2025 Do: 8.10.2026	Unieważniony przez admin	
Główny prompt dla konwersacyjnego Agenta AI. 622050af-cbab-4440-9ba4-d6cc5333caf9	Właściciel: admin EnterSoft Utworzył: admin	0 użyc	Nigdy	Od: 8.10.2025 Do: 8.10.2026	Unieważniony przez System	
Integracja z ThunderBird 03777ad4-3435-4a19-a87a-e8a23923c85a	Właściciel: admin EnterSoft Utworzył: admin	0 użyc	Nigdy	Od: 8.10.2025 Do: 8.10.2026	Unieważniony przez System	

Główna tabela w panelu przedstawia listę wszystkich tokenów i jest centrum audytu oraz zarządzania. Każda kolumna dostarcza kluczowych informacji o cyklu życia i użyciu tokenów:

- **Token:**
 - **Opis:** Czytelna nazwa tokena ustawiona podczas jego tworzenia.
 - **JTI:** Unikalny identyfikator tokena w formacie UUID, który jest jego faktycznym kluczem w bazie danych.
- **Właściciel / Utworzony przez:**
 - **Właściciel:** Wskazuje, z którym kontem użytkownika powiązany jest token. Aplikacja zewnętrzna używająca tego tokena działa z uprawnieniami tego użytkownika.
 - **Utworzył:** Pokazuje, który administrator wygenerował dany token.
- **Statystyki Użycia:**
 - **Liczba użyc:** Łączna liczba pomyślnych wywołań API przy użyciu tego tokena.
 - **Błędy:** Liczba zanotowanych nieudanych prób autoryzacji (np. z powodu użycia wygasłego lub unieważnionego tokena). Jest to ważny wskaźnik bezpieczeństwa.
- **Ostatnia Aktywność:**

- **Data i godzina** ostatniego pomyślnego użycia tokena.
- **Adres IP**, z którego nastąpiło ostatnie połączenie.
- **Ważność:** Precyzyjne daty utworzenia i wygaśnięcia tokena.
- **Status:** Dynamicznie określany stan tokena:
 - **Aktywny:** Token jest ważny i gotowy do użycia.
 - **Nieaktywny:** Token jest wciąż ważny, ale nie był używany przez ostatnie 30 dni.
 - **Wygaś:** Data ważności tokena minęła.
 - **Unieważniony:** Token został ręcznie odwołany przez administratora.
- **Akcje:** Zestaw ikon umożliwiających szybkie zarządzanie każdym tokenem:
 -  **(Pokaż logi użycia):** Otwiera okno ze szczegółową historią ostatnich 100 operacji wykonanych przy użyciu danego tokena. W logach znajdują się: dokładny czas, wywołany endpoint, kod odpowiedzi HTTP, czas trwania operacji oraz adres IP. Dodatkowo, w oknie tym wyświetlane jest podsumowanie statystyk wydajności (średni czas odpowiedzi, liczba błędów) **tylko dla tego tokena**.

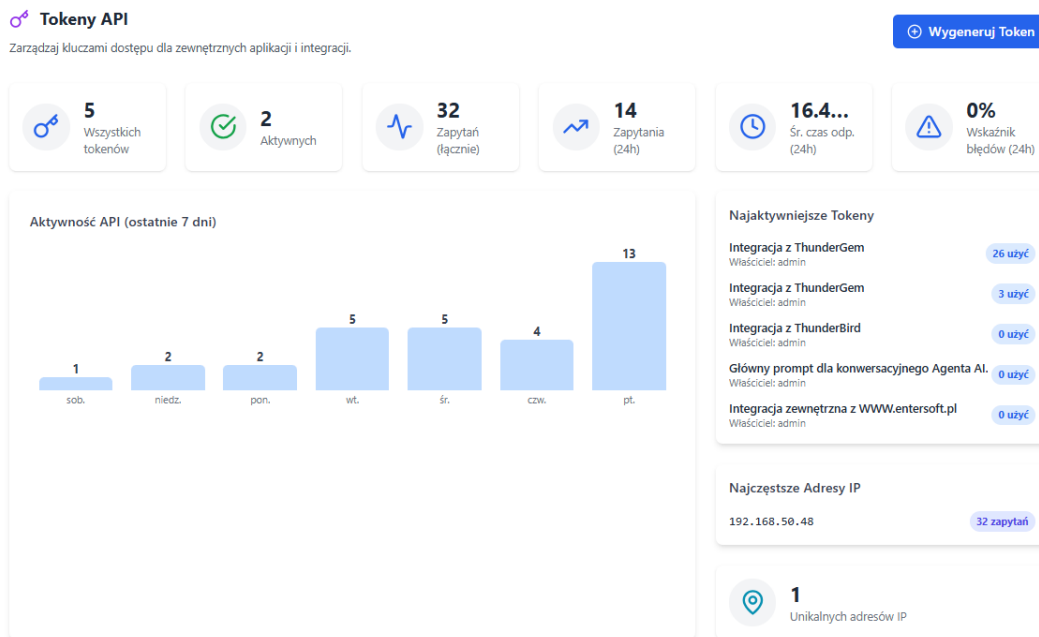
Logi Użycia Tokena: Integracja z ThunderGem ×

3 Wszystkich zapytań	42.00 ms Śr. czas odpowiedzi	0 Zapytań z błędem
--------------------------------	--	------------------------------

Czas ▼	Endpoint	Status	Czas trwania (ms)	Adres IP
10.10.2025, 07:08:28	/api/v1/conversations/rag-stream	200	9	192.168.50.
10.10.2025, 07:08:23	/auth/models	200	89	192.168.50.
10.10.2025, 07:08:23	/api/v1/meta/agents/list	200	28	192.168.50.

-  **(Odnów token):** Pozwala na bezpieczną regenerację klucza. Ta akcja natychmiast unieważnia stary token i generuje nowy z identycznymi uprawnieniami, opisem i datą ważności. Nowy token należy skopiować, ponieważ nie będzie on ponownie wyświetlony.
-  **(Unieważnij token):** Natychmiastowo i trwale odwołuje token, uniemożliwiając jego dalsze użycie. Wszystkie kolejne próby autoryzacji z tym tokenem zostaną odrzucone i zalogowane jako błąd. System zapisze, który administrator dokonał unieważnienia.

Panel Statystyk



Na górze panelu znajduje się panel prezentujący kluczowe wskaźniki (KPI) dotyczące wykorzystania API:

- **Liczba Tokenów:** Całkowita i aktywna liczba wygenerowanych tokenów.
- **Liczba Zapytań:** Łączna liczba zapytań obsłużonych przez tokeny oraz liczba zapytań w ciągu ostatnich 24 godzin.
- **Aktywność API:** Wykres słupkowy obrazujący liczbę wywołań API w podziale na ostatnie 7 dni.
- **Najaktywniejsze Tokeny:** Ranking 5 tokenów z największą liczbą użyci, wraz z informacją o ich właścicielu i przeznaczeniu.
- **Średni Czas Odpowiedzi:** Wskazuje średni czas (w milisekundach) potrzebny na przetworzenie zapytania API w ciągu ostatnich 24 godzin. Jest to kluczowy wskaźnik wydajności systemu. Wartości powyżej **500 ms** mogą być sygnalizowane kolorem ostrzegawczym, wskazując na potencjalne spowolnienia.
- **Wskaźnik Błędów:** Pokazuje procent zapytań z ostatnich 24 godzin, które zakończyły się błędem (np. z powodu błędnych danych wejściowych lub problemów serwera). Jest to najważniejszy wskaźnik stabilności integracji. Wartości powyżej **5%** są sygnalizowane na czerwono.
- **Najczęstsze Adresy IP:** Ranking pięciu adresów IP, z których najczęściej korzystano z tokenów API. Umożliwia to monitorowanie, czy zapytania pochodzą z oczekiwanych i zaufanych źródeł.
- **Unikalne Adresy IP:** Karta pokazująca łączną liczbę unikalnych adresów IP, które użyły tokenów. Nagły wzrost tej wartości może sygnalizować nową, nieprzewidzianą integrację lub próbę nieautoryzowanego dostępu.

Filtrowanie i Wyszukiwanie Tokenów

🔍 Szukaj po opisie lub JTI...

Aby ułatwić zarządzanie dużą liczbą kluczy API, nad główną tabelą wprowadzono zaawansowane opcje filtrowania:

- **Szukaj po opisie lub JTI:** Pole tekstowe umożliwia dynamiczne przeszukiwanie listy tokenów. Wystarczy wpisać fragment opisu lub unikalnego identyfikatora (JTI), aby natychmiast zawęzić wyniki.
- **Filtruj wg Statusu:** Rozwijana lista pozwala wyświetlić tylko te tokeny, które spełniają określone kryterium:
 - **Aktywne:** Ważne i nieunieważnione.
 - **Wygaste:** Tokeny, których data ważności minęła.
 - **Unieważnione:** Tokeny ręcznie odwołane przez administratora.
- **Filtruj wg Organizacji:** (Dostępne dla Super Administratora) Pozwala na wyświetlanie tokenów należących tylko do wybranej organizacji.

Generowanie Nowego Tokena

Wygeneruj Nowy Token Serwisowy

Organizacja

EnterSoft

Użytkownik (w imieniu którego będzie działał token)

admin (SuperAdmin)

Opis

np. Integracja z systemem CRM

Ważność (dni)

365

Uprawnienia (Scopes)

Szablon: Dostęp do AI

Szablon: Zarządzanie Dokumentami

☒ AI Fakt: Pytania

☒ AI Agent: Pełny dostęp

☒ Dokumenty: Odczyt

☐ Dokumenty: Zapis

Anuluj

Generuj

Aby utworzyć nowy token, należy kliknąć przycisk "**Wygeneruj Token**" i wypełnić formularz:

- **Organizacja:** (Tylko Super Administrator) Wybór organizacji, w kontekście której będzie działał token.
- **Użytkownik:** Wskazanie użytkownika, w imieniu którego aplikacja zewnętrzna będzie się uwierzytelniać. Token odziedziczy jego uprawnienia.
- **Opis:** Krótka, zrozumiała nazwa pozwalająca zidentyfikować przeznaczenie tokena (np. "Integracja z ...").
- **Ważność (dni):** Okres aktywności tokena (domyślnie 365 dni).

- **Uprawnienia:** Precyzyjny wybór uprawnień, które chcesz nadać tokenowi. Dostępne są tylko te uprawnienia, które posiada wybrany użytkownik, co uniemożliwia eskalację uprawnień. Można skorzystać z gotowych szablonów, np. "Dostęp do AI".

Po wygenerowaniu token pojawi się w oknie modalnym. **Należy go natychmiast skopiować i bezpiecznie zapisać, ponieważ nie będzie możliwości jego ponownego odtworzenia.**

Token Wygenerowany Pomyślnie

Skopiuuj token teraz! Nie będzie on widoczny ponownie.

```
eyJhbGciOiJIUzI1NiIsInR5cCI6IkpXVCJ9.eyJzdWIiOiIxIiwidXNlcm5hbWUiOiJhZG1pbiIsInJvbGUiOiJTdXB1ckFkbWluIiwib
```



Rozumiem, zamknij

Przeglądanie i Monitorowanie Tokenów

Główna tabela w panelu przedstawia listę wszystkich tokenów wraz ze szczegółowymi informacjami audytowymi:

- **Opis / Status:** Identyfikator tokena oraz jego aktualny stan (np. Aktywny, Wygaśł, Nieaktywny, Unieważniony).
- **Właściciel / Utworzony przez:** Wskazuje, z którym kontem powiązany jest token i który administrator go utworzył.
- **Statystyki Użycia:** Łączna liczba pomyślnych użyci oraz liczba zanotowanych błędów autoryzacji
- **Ostatnia Aktywność:** Data i godzina ostatniego użycia tokena oraz adres IP, z którego nastąpiło połączenie.
- **Szczegóły Czasowe:** Precyzyjne daty utworzenia i wygaśnięcia.

Unieważnianie Tokena (Revoke)

Jeśli token został skompromitowany lub integracja przestała być używana, Super Administrator może go w każdej chwili unieważnić:

1. Kliknij ikonę  w wierszu wybranego tokena.
2. Potwierdź operację.

Token zostanie natychmiast oznaczony w bazie danych jako unieważniony, a wszystkie kolejne próby jego użycia zostaną odrzucone. System zapisze również, który administrator dokonał unieważnienia.

Zarządzanie Modelami AI

Zakładka **Zarządzanie Modelami AI** to Twoje centrum kontroli nad "mózgiem" systemu – Dużymi Modelami Językowymi (LLM), które generują odpowiedzi. Panel ten pozwala na dynamiczne dodawanie nowych modeli, przeglądanie już zainstalowanych oraz usuwanie nieużywanych w celu optymalizacji zasobów serwera. Elastyczność w zarządzaniu modelami pozwala dostosować system do specyficznych potrzeb, np. używając modeli wyspecjalizowanych w analizie kodu, tekstów prawniczych czy konwersacji.

Główne Elementy Panelu

Pobierz nowy model

Wpisz nazwę modelu z [biblioteki Ollama](#) (np. 'mistral:7b') lub wybierz z listy rekomendowanych.

Pobierz

Zainstalowane modele

NAZWA	ROZMIAR	MODYFIKACJA	AKCJE
deepseek-r1:latest	4.87 GB	1.10.2025, 09:13:15	
llama3:latest	4.34 GB	1.10.2025, 09:01:21	
llama3:8b	4.34 GB	1.10.2025, 09:00:35	
gemma:7b	4.67 GB	1.10.2025, 08:58:01	
phi3:latest	2.03 GB	1.10.2025, 08:18:40	
gemma:2b	1.56 GB	1.10.2025, 08:16:54	
mistral:latest	4.07 GB	1.10.2025, 08:16:04	

llama3:latest

llama3:8b

llama3:70b

mistral:latest

mistral:7b-instruct-v0.2-q4_k_s

mistral-nemo:latest

gemma:2b

gemma:7b

phi3:latest

phi3:medium

magistral:latest

nous-hermes2:latest

deepseek-llm:latest

openchat:latest

openhermes:latest

Interfejs składa się z trzech logicznych części, które pozwalają na pełne zarządzanie cyklem życia modeli w Twoim systemie.

1. Pobierz Nowy Model Ta sekcja umożliwia rozszerzenie możliwości systemu o nowe modele z oficjalnej [biblioteki Ollama](#).

- **Pole tekstowe:** Wpisz dokładną nazwę modelu, który chcesz pobrać (np. llama3:8b lub mistral:7b-instruct-v0.2-q4_k_s).
- **Przycisk Pobierz:** Rozpoczyna proces ściągania modelu na serwer. Operacja odbywa się w tle, a jej postęp jest widoczny na pasku postępu, który informuje o aktualnym statusie (np. "ściąganie manifestu", "weryfikacja", "rozpakowywanie").

2. Rekomendowane Modele Aby ułatwić wybór, po prawej stronie znajduje się lista modeli rekomendowanych i przetestowanych przez zespół AsWiseAI.

- Kliknięcie ikony + obok nazwy modelu automatycznie wypełni pole "Pobierz nowy model", co upraszcza i przyspiesza proces dodawania sprawdzonych LLM-ów.

3. Zainstalowane Modele Główna część panelu to lista modeli, które są obecnie dostępne dla AsWiseAI. Tabela zawiera kluczowe informacje:

- **Nazwa:** Pełna nazwa modelu wraz z tagiem (np. llama3:latest). Kliknięcie nazwy przeniesie Cię do oficjalnej strony modelu w bibliotece Ollama.
- **Rozmiar:** Ilość miejsca na dysku zajmowana przez dany model.
- **Modyfikacja:** Data ostatniej modyfikacji lub pobrania modelu.
- **Akcje:** Umożliwia trwałe usunięcie modelu z serwera za pomocą ikony kosza, co zwalnia przestrzeń dyskową.

Wskazówki i Dobre Praktyki

- **Rozmiar i Wydajność:** Nazwy modeli często zawierają informację o ich rozmiarze (np. 7b, 13b, 70b - miliardy parametrów). Większe modele są zazwyczaj bardziej zaawansowane i generują odpowiedzi wyższej jakości, ale wymagają znacznie więcej pamięci VRAM i przestrzeni dyskowej.
- **Kwantyzacja (Jakość vs. Zasoby):** Zwróć uwagę na dodatkowe tagi w nazwie, takie jak q4_K_M czy Q8_0. Określają one poziom kwantyzacji (kompresji) modelu. Modele o niższej kwantyzacji (np. Q4) zajmują mniej miejsca i VRAM, ale mogą oferować nieco niższą precyzję niż ich większe odpowiedniki (np. Q8).
- **Zarządzanie Zasobami:** Regularnie usuwaj modele, których już nie używasz. Modele językowe potrafią zajmować od kilku do kilkudziesięciu gigabajtów przestrzeni dyskowej, więc utrzymywanie porządku jest kluczowe dla stabilności serwera.

Zarządzanie regułami czyszczenia

System AsWiseAI został wyposażony w mechanizm automatycznego oczyszczania zaśmieconych odpowiedzi, który jest dostępny z pozycji USTAWIENIA -> Reguły (Few-Shot)

Ustawienia

Konto

Historia

Reguły (Few-Shot)

Zarządzanie Regułami Czyszczenia AI

Importuj

Eksportuj

Dodaj Regułę

Zarządzaj regułami (wyrażeniami regularnymi), które automatycznie czyszczą odpowiedzi modeli językowych.

Typ	Wzorzec (Regex)	Opis	Status	Akcje
prefix	^(Pytanie\s*\d*)?\s*[:\"]*\s*	Usuwa "Pytanie 2:", "Pytanie:", ":" etc.	ON	
prefix	^(?:Here are Oto Wygenerowane)\s*.*?\s*	Usuwa "Here are the 3 questions:" etc.	ON	
prefix	^\d*.*\s*	Usuwa punktory i numerację, np. "1.", "2.", "3."	OFF	
suffix	\s*(Translation:.*?)\$	Usuwa "(Translation: ...)" na końcu	ON	
suffix	\s*(What are.*?)\$	Usuwa "(What are...)" na końcu	ON	
prefix	^(Here is the answer in HTML format\s*\d*)?\s*[:\"]*\s*	[Kopia] Usuwa "Here is the answer in HTML format..."	ON	

- Na stronie Reguły (Few-Shot) (/admin/cleaning-patterns) administratorzy mogą definiować reguły (w postaci wyrażeń regularnych), które automatycznie usuwają niechciane fragmenty z odpowiedzi AI.
- Reguły te mogą być aktywowane, dezaktywowane, edytowane lub usuwane.
- Podobnie jak w przypadku promptów, dostępny jest import i eksport reguł w formacie JSON.

Kopia zapasowa i resetowanie organizacji

Zarządzanie Kopiami Zapasowymi

Kopie zapasowe
 ➕ Utwórz kopię 🗑️ Zachowaj tylko 3

NAZWA PLIKU	UTWORZONO	ROZMIAR	AKCJE
qdrant_backup_20250812_171304.zip	12.08.2025, 17:13:15	213.01 MB	Przywróć Pobierz Usuń
qdrant_backup_20250811_085425.zip	11.08.2025, 08:54:34	209.19 MB	Przywróć Pobierz Usuń
qdrant_backup_20250807_075931.zip	7.08.2025, 07:59:38	142.41 MB	Przywróć Pobierz Usuń
qdrant_backup_20250806_215223.zip	6.08.2025, 21:52:26	65.85 MB	Przywróć Pobierz Usuń
qdrant_backup_20250805_160550.zip	5.08.2025, 16:05:53	68.72 MB	Przywróć Pobierz Usuń

- Strona Backupy (/kopie) oferuje narzędzia do tworzenia i zarządzania kopiami zapasowymi całego systemu (Qdrant, bazy SQL i wgranych plików).
- **Tworzenie kopii:** Po kliknięciu Utwórz kopię, system tworzy archiwum ZIP z kompletnym stanem aplikacji.
- **Przywracanie:** Można przywrócić system z dowolnego pliku kopii zapasowej.

Resetowanie danych

- **Globalny reset:** Na stronie Reset bazy (/reset_all), po podaniu specjalnego PIN-u, możliwe jest trwałe i nieodwracalne usunięcie **wszystkich danych dla wszystkich organizacji** z systemu i odtworzenie go do stanu początkowego.

Zarządzanie danymi Organizacji

Wybierz akcję, którą chcesz wykonać dla danej organizacji.

EnterSoft (ID: 1)
 [Resetuj](#)
[Usuń](#)

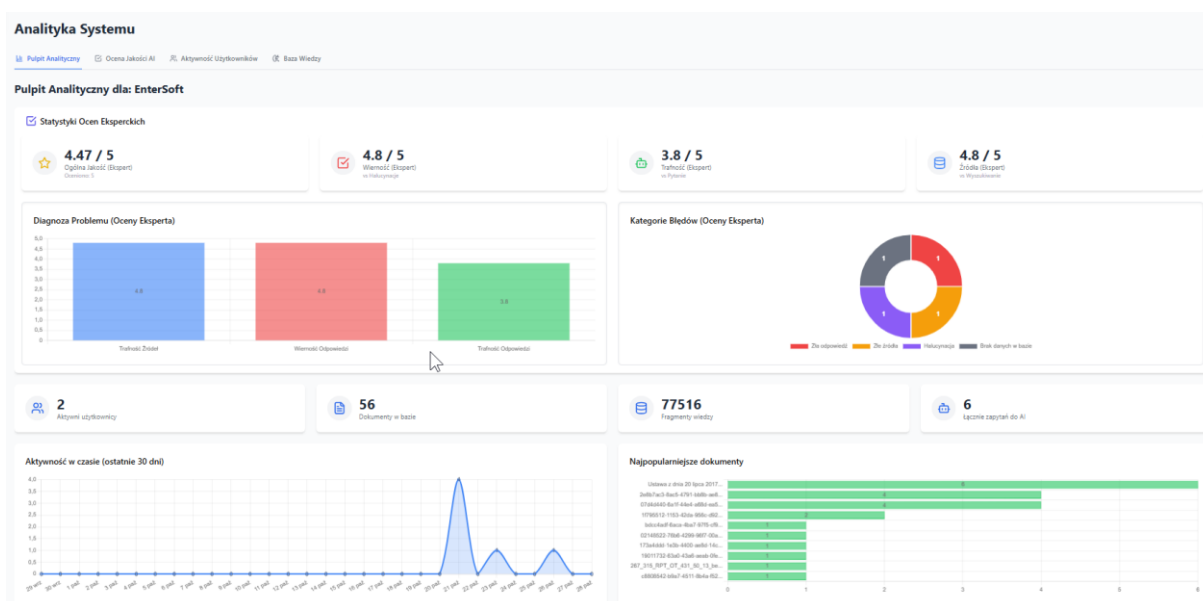
⚠️ Globalny Reset Systemu

Ta operacja bezpowrotnie usunie ****WSZYSTKIE**** dane ze wszystkich organizacji, wszystkich użytkowników, wszystkie dokumenty i wszystkie kolekcje Qdrant. Używaj z najwyższą ostrożnością.

Wprowadź globalny PIN administratora

RESETUJ CAŁY SYSTEM

8. Panel Analityczny



Panel Analityczny to centrum monitorowania dostępne dla administratorów systemu. Zapewnia on wgląd w kluczowe wskaźniki wykorzystania aplikacji, aktywność użytkowników oraz, co najważniejsze, w jakość odpowiedzi dostarczanych przez AI.

Panel dostępny jest z bocznego menu nawigacyjnego po kliknięciu **Analitika**.

Panel Analityczny podzielony jest na cztery główne zakładki:

1. **Pulpit Analityczny** (Widok główny)
2. **Ocena Jakości AI** (Moduł ewaluacji RAG)
3. **Aktywność Użytkowników**
4. **Baza Wiedzy**

1. Pulpit Analityczny

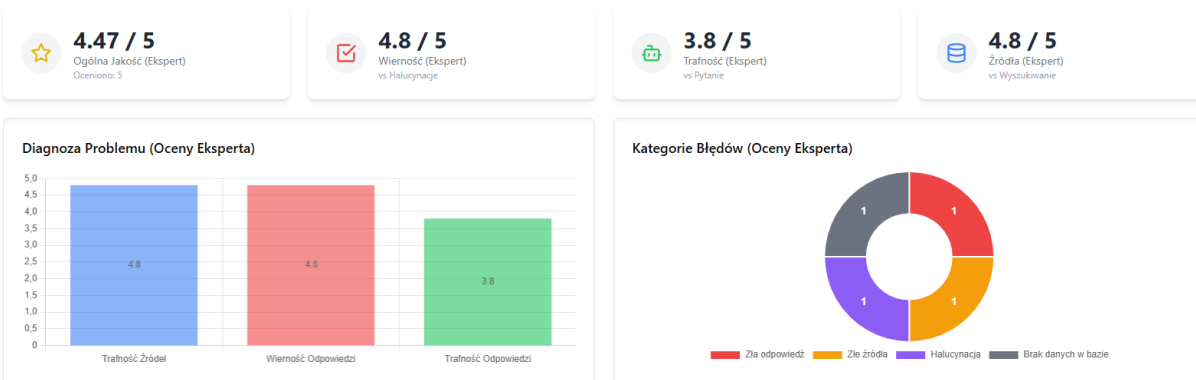
Główny pulpit prezentuje zagregowane dane dotyczące całej organizacji.

Kluczowe Wskaźniki Efektywności (KPI): W górnej części panelu znajdują się cztery główne karty KPI:

- **Aktywni użytkownicy:** Łączna liczba kont użytkowników w organizacji.
- **Dokumenty w bazie:** Całkowita liczba unikalnych dokumentów (plików) wgranych do systemu.
- **Fragmenty wiedzy:** Łączna liczba wektorów (fragmentów tekstu) w bazie wiedzy.
- **Łącznie zapytań do AI:** Suma wszystkich zapytań (pytań) zadanych przez użytkowników w trybie "AI Fakt".

Statystyki Ocen Eksperskich

Statystyki Ocen Eksperskich



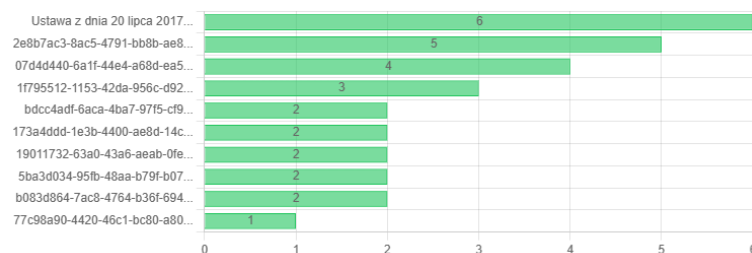
Ta sekcja jest bezpośrednio powiązana z modułem "Ocena Jakości AI". Pokazuje ona:

- **Średnie oceny (w skali 1-5):** Pokazuje zagregowaną, średnią ocenę wystawioną przez ekspertów w trzech kluczowych kategoriach:
 - **Ogólna Jakość** (Czy AI znalazło właściwe dokumenty?)
 - **Wierność Odpowiedzi** (Czy AI nie halucynowało i trzymało się faktów?)
 - **Trafność Odpowiedzi** (Czy AI odpowiedziało na zadane pytanie?)
- **Kategorie Błędów:** Wykres kołowy pokazujący najczęstsze przyczyny błędnych odpowiedzi (np. "Halucynacja", "Złe źródła") zidentyfikowane przez ekspertów.

Wykresy Aktywności:

- **Aktywność w czasie (ostatnie 30 dni):** Wykres liniowy pokazujący trend liczby zapytań do AI w ujęciu dziennym.
- **Najpopularniejsze dokumenty:** Wykres słupkowy pokazujący, które dokumenty są najczęściej cytowane (używane jako źródło) w odpowiedziach AI.

Najpopularniejsze dokumenty



2. Ocena Jakości AI (Ewaluacja RAG)

Jest to kluczowy moduł dla administratorów i ekspertów merytorycznych, służący do monitorowania i „ręcznej oceny” (Human Evaluation) jakości odpowiedzi generowanych przez system AI.

1. Zakładka "Do Oceny"

Pytanie Użytkownika

Czy istnieją efekty uboczne związane z wykorzystaniem substancji szkodliwych w rolnictwie?

Odpowiedź AI (Widok Użytkownika)

Tak, istnieją efekty uboczne związane z wykorzystaniem substancji szkodliwych w rolnictwie. Na podstawie dostarczonych dokumentów [1-15], mogę stwierdzić, że stosowanie środków chemicznych w uprawach roślin zbożowych może przyczynić się do zwiększenia szkodliwości [2]. Warto również uważać na toksykologię substancji i ich wpływ na rośliny uprawne i środowisko [3].

Warto pamiętać, że herbicydy stosowane w przedplonach mogą stanowić zagrożenie dla zdrowia człowieka i zwierząt [5]. Ponadto, uszkodzenia powodowane przez herbicydy mogą w różny sposób wpływać na rośliny uprawne [7].

Warto również pamiętać, że zabiegi agrotechniczne, biologia i fitotoksyczność na uprawianych następnie roślinach są ważnymi czynnikami, które mogą wpłynąć niekorzystnie na plon i jego jakość [12].

► Źródła (5)

🔍 Pokaż Kontekst (Źródła) dostarczone do AI

<> Pokaż Odpowiedź AI (Surowy Tekst)

Ta zakładka prezentuje kolejną zapytań i odpowiedzi, które nie zostały jeszcze zweryfikowane przez eksperta.

- **Interfejs Oceny:** Administrator widzi dokładne pytanie użytkownika oraz odpowiedź AI sformatowaną w ten sam sposób, w jaki widział ją użytkownik (wraz z linkami do źródeł).
- **Panele Diagnostyczne:** Dostępne są zwijane panele "Pokaż Kontekst" (pokazuje fragmenty, które AI otrzymało jako bazę do odpowiedzi) oraz "Pokaż Surowy Tekst" (czysta odpowiedź modelu).
- **Formularz Oceny:** Ekspert ocenia odpowiedź za pomocą 5-gwiazdkowego systemu w trzech kluczowych dla systemów RAG kategoriach:

Formularz Oceny

1. Ocena Wyszukiwania (Context Relevance)

☆☆☆☆☆ 0/5

2. Ocena Wierności (Faithfulness)

☆☆☆☆☆ 0/5

3. Ocena Trafności (Answer Relevance)

☆☆☆☆☆ 0/5

Kategoria Błędu (Opcjonalnie)

Brak (ocena pozytywna)

Notatki (Opcjonalnie)

Dodatkowe uwagi dla zespołu deweloperskiego...

Pozostało w kolejce: 2

⏮️ Pomiń

🚀 Zatwierdź i weź następne

1. **Ocena Wyszukiwania:** Czy AI znalazło właściwe fragmenty?
2. **Ocena Wierności:** Czy AI nie "zmyślało" i trzymało się faktów ze znalezionych źródeł?
3. **Ocena Trafności:** Czy AI odpowiedziało na temat?

- **Kategoryzacja Błędów**

Kategoria Błędu (Opcjonalnie)

Brak (ocena pozytywna)	▼
Brak (ocena pozytywna)	
Błąd: Złe źródła (AI nie znalazło dobrych dokumentów)	
Błąd: Halucynacja (AI zmyślało, mimo dobrych źródeł)	
Błąd: Zła odpowiedź (AI odpowiedziało nie na temat)	
Informacja: Brak danych w bazie (AI słusznie nic nie znalazło)	

W przypadku złej odpowiedzi, ekspert może wskazać główną przyczynę błędu (np. Halucynacja, Złe źródła).

- **Notatki:** Możliwość dodania opcjonalnego komentarza.

Naciśnięcie "Zatwierdź" zapisuje ocenę i automatycznie łąduje kolejny element z kolejki.

2. Zakładka "Ocenione"

Do Oceny		Ocenione					
PYTANIE I NOTATKI		OCENA OGÓLNA	ŹRÓDŁA	WIERNOŚĆ	TRAFNOŚĆ	KATEGORIA BŁĘDU	AKCJE
W jakich sytuacjach nasilenie się konwergencji zjawisk w przyrodzie prowadzi do wyjątkowych rozwiązań? Oceniono przez: admin w dniu: 28.10.2025		4.67/ 5	5	5	4	Brak danych w bazie	 
W jakich sytuacjach nasilenie się konwergencji zjawisk w przyrodzie prowadzi do wyjątkowych rozwiązań? Oceniono przez: admin w dniu: 28.10.2025		4.00/ 5	4	4	4	Zła odpowiedź	 
W jakim stopniu cechy środowiska wód morskich wpływają na życie organizmów? Oceniono przez: admin w dniu: 28.10.2025		4.00/ 5	5	5	2	Halucynacja	 
W jaki sposób technologia wpływa na nasze prawa i wolności? Oceniono przez: admin w dniu: 28.10.2025		5.00/ 5	5	5	5	Złe źródła	 
Jaka jest maksymalna kwota dotacji z programu Czyste powietrze? Przedstaw w formie tabeli wszystkie opcje dofinansowania Oceniono przez: admin w dniu: 28.10.2025		4.67/ 5	5	5	4	Brak	 

Razem: 5 (Strona 1 z 1)

<

>

Wyświetla listę wszystkich historycznych ocen. Administrator może tu przeglądać średnie wyniki dla każdej odpowiedzi, edytować lub usuwać poprzednie oceny. Usunięcie oceny powoduje jej powrót do kolejki "Do Oceny".

3. Aktywność Użytkowników

Ta zakładka koncentruje się na analizie zachowań użytkowników w organizacji.

- Najaktywniejsi Użytkownicy:** Tabela rankingowa użytkowników posortowana według łącznej liczby zapytań zadanych do AI.

Analityka Systemu	
Pulpit Analityczny	Ocena Jakości AI
Aktywność Użytkowników	Baza Wiedzy
Najaktywniejsi Użytkownicy	
UŻYTKOWNIK	LICZBA ZAPYTAŃ
jacko jacko@entersoft.pl	18
Użytkownicy Wymagający Uwagi	
Aktywni użytkownicy, którzy nie logowali się od ponad 30 dni lub nigdy.	
UŻYTKOWNIK	OSTATNIE LOGOWANIE
Wszyscy użytkownicy są aktywni.	

- Użytkownicy Wymagający Uwagi:** Lista aktywnych kont, które nie logowały się do systemu od ponad 30 dni.

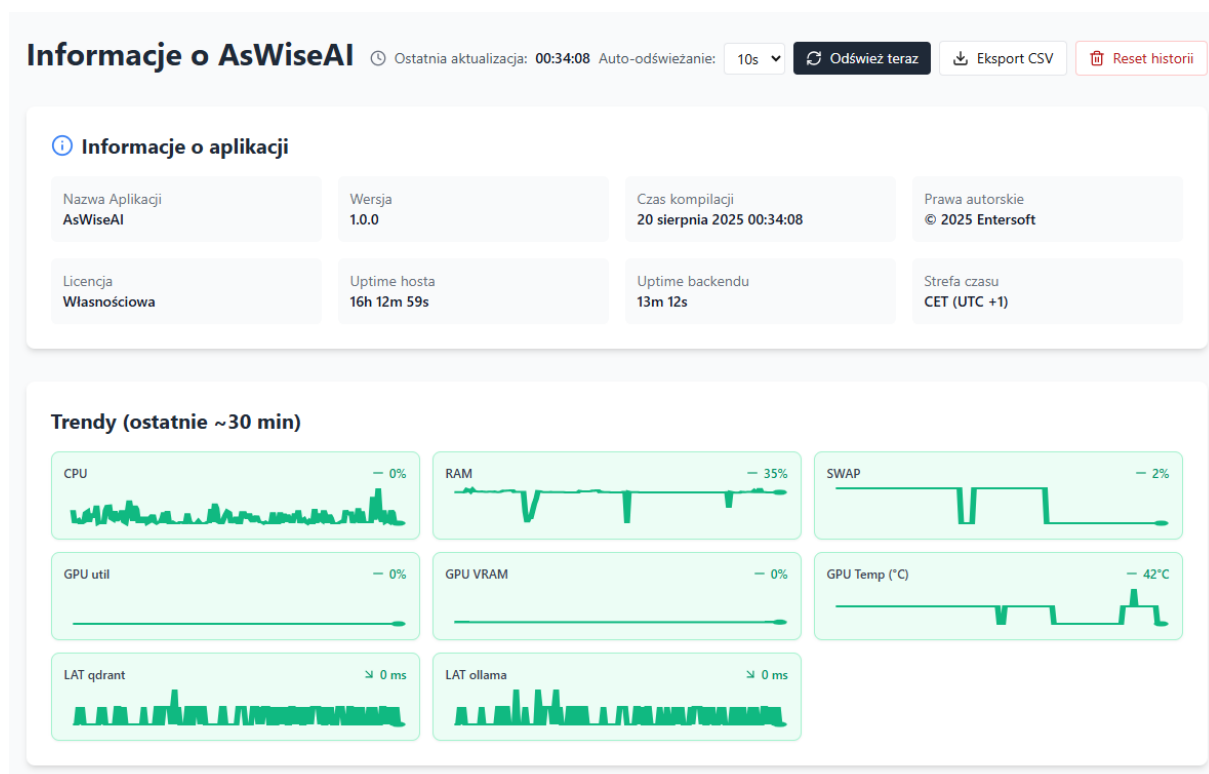
4. Baza Wiedzy

Pulpit Analityczny Ocena Jakości AI Aktywność Użytkowników Baza Wiedzy	
Najczęściej Wykorzystywane Dokumenty	
DOKUMENT	LICZBA CYTOWAŃ
CPZ_Podrecznik_1.pdf	15
PROGRAM PRIORYTETOWY Tytuł programu- Czyste Powietrze_00_ppcp_z_czask_fm_zal1_do_uchwaly.pdf	12
Instrukcja wypełniania wniosku o płatność.pdf	11
Program priorytetowy Czyste Powietrze - Część 4i.pdf	11
PROGRAM PRIORYTETOWY Tytuł programu- Czyste Powietrze c.3_instrukcja_wsp_suf_s_b_.pdf	10
zakres wydatkowania środków Na co i ile - zakres i wysokość dofinansowania w programie Czyste Powietrze.pdf	10
Regulamin naboru wniosków o dofinansowanie - Część 4i PPCh.pdf	8
PROGRAM PRIORYTETOWY Tytuł programu- Czyste Powietrze a.8_instrukcja_dpa_suf_s_b_.pdf	8
PROGRAM PRIORYTETOWY Tytuł programu- Czyste Powietrze a.1_regulamin_ppcp_suf_.pdf	7
PROGRAM PRIORYTETOWY Tytuł programu- Czyste Powietrze a.16_instrukcja_dpr_suf_.pdf	7
Zapytania Bez Znalazionych Źródeł	
Lista ostatnich zapytań, które wskazują na potencjalne luki w bazie wiedzy.	
PYTANIE	UŻYTKOWNIK
Brak zapytań bez znalezionych źródeł.	

Ta zakładka analizuje stan samej bazy wiedzy i pomaga identyfikować luki informacyjne.

- **Najczęściej Wykorzystywane Dokumenty:** Tabela pokazująca, które dokumenty są najczęściej używane przez AI do udzielania odpowiedzi.
- **Zapytania Bez Znalazionych Źródeł:** Niezwykle ważna tabela, która listuje pytania użytkowników, na które system AI nie był w stanie znaleźć żadnych pasujących fragmentów w bazie wiedzy. Pozwala to administratorom szybko zidentyfikować brakujące informacje i uzupełnić bazę o odpowiednie dokumenty.

9. Status systemu



1 Co to za strona i do czego służy

Status systemu to kokpit operacyjny AsWiseAI. W jednym miejscu pokazuje:

- kondycję kluczowych usług (np. **Ollama**, **Qdrant**),
- bieżące metryki serwera (CPU, RAM, SWAP, dyski),
- informacje o **GPU** (jeśli jest dostępne),
- czas odpowiedzi (latencję) wybranych integracji,
- krótkoterminowe **trendy** i proste alerty ostrzegające o przeciążeniach.

Dzięki temu szybko sprawdzisz, czy system działa poprawnie.

2 Dostęp i nawigacja

- **Ścieżka:** Menu → Status systemu (URL: `/system-status`).
- **Upewnienie wymagane ADMINISTRATORA:** „Podgląd Stanu Systemu”.

Jeśli go nie masz, poproś Administratora lub SuperAdministratora.

3 Jak czytać ekran – układ i elementy

Pasek sterowania (u góry po prawej)

- **Ostatnia aktualizacja** – pokazuje czas ostatniego pobrania danych.

- **Auto-odświeżanie:** wybierz Off / 10s / 30s / 60s.
Ustawienie jest zapamiętywane w przeglądarce i działa, dopóki masz otwartą stronę.
- **Odśwież teraz** – natychmiast pobiera świeże dane.
- **Eksport CSV** – zapisuje do pliku historię trendów z Twojej przeglądarki (przydatne do analizy/zgłoszeń).
- **Reset historii** – czyści lokalną historię trendów (działa tylko w Twojej przeglądarce).

Baner alertów (jeśli widoczny)

Proste ostrzeżenia oparte o **średnią z ostatniej minuty** (przy interwale 10 s).

Kolory:

- **Zielony** – w normie.
- **Bursztynowy** – ostrzeżenie (podwyższone wartości).
- **Czerwony** – krytyczne (trwale wysokie wartości).

Dotyczy: CPU, RAM, SWAP, GPU (użycie/VRAM/temperatura) oraz **latencji** integracji.

Informacje o aplikacji

- **Nazwa, Wersja, Czas kompilacji** (znacznik czasu backendu),
- **Prawa autorskie, Licencja**,
- **Uptime hosta** (jak długo działa system operacyjny),
- **Uptime backendu** (jak długo działa proces backendu),
- **Strefa czasu serwera**.

Status usług

Dwie karty: **Serwer Ollama** i **Baza wektorowa Qdrant**.

Status usług

 Serwer Ollama	
Wersja:	0.10.1
Endpoint:	http://ollama:11434
Modele (liczba):	3
▶ Szczegóły	

 Baza wektorowa Qdrant	
Wersja:	N/A
Alive:	Nie
Endpoint:	http://qdrant:6333
Kolekcje (liczba):	2
▶ Szczegóły	

Każda karta pokazuje:

- **Lampkę stanu** (zielona/czerwona),
- **Wersję**,
- **Alive** (czy odpowiada), **Endpoint**,
- Liczbę **modeli** (Ollama) lub **kolekcji** (Qdrant),

- **Szczegóły** (po rozwinięciu): podgląd list, podstawowa telemetria i **surowe dane** (JSON).

Wskazówka: jeśli lampka jest czerwona albo widzisz błąd, sprawdź szczegóły – często podają przyczynę (np. „connection refused”).

Trendy (mini-wykresy)

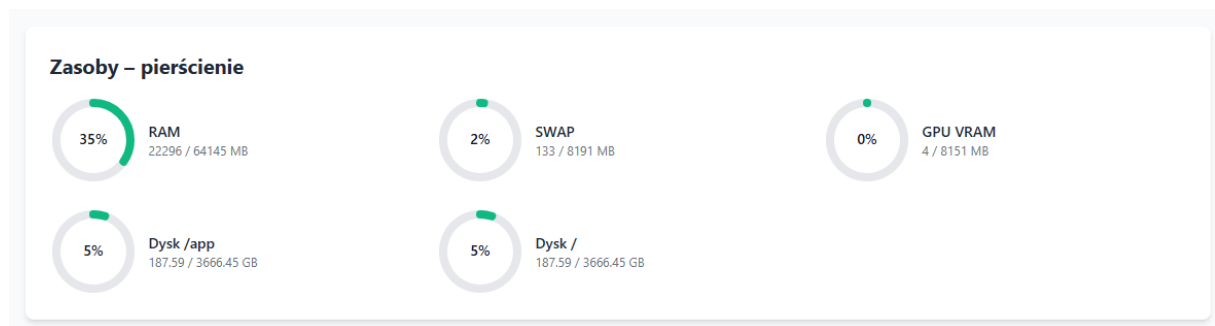
Krótko-okresowe, lekkie wykresy wartości (ostatnie ~30 minut):

- **CPU, RAM, SWAP**, (jeśli dostępne) **GPU util, GPU VRAM, GPU temperatura**,
- **LAT** (latencja) dla kilku integracji.

U góry wykresu: ostatnia wartość i strzałka kierunku (wzrost/spadek/brak zmiany).
Dane do trendów są **trzymane lokalnie w przeglądarce** (nie na serwerze).

Zasoby – pierścienie (donuty)

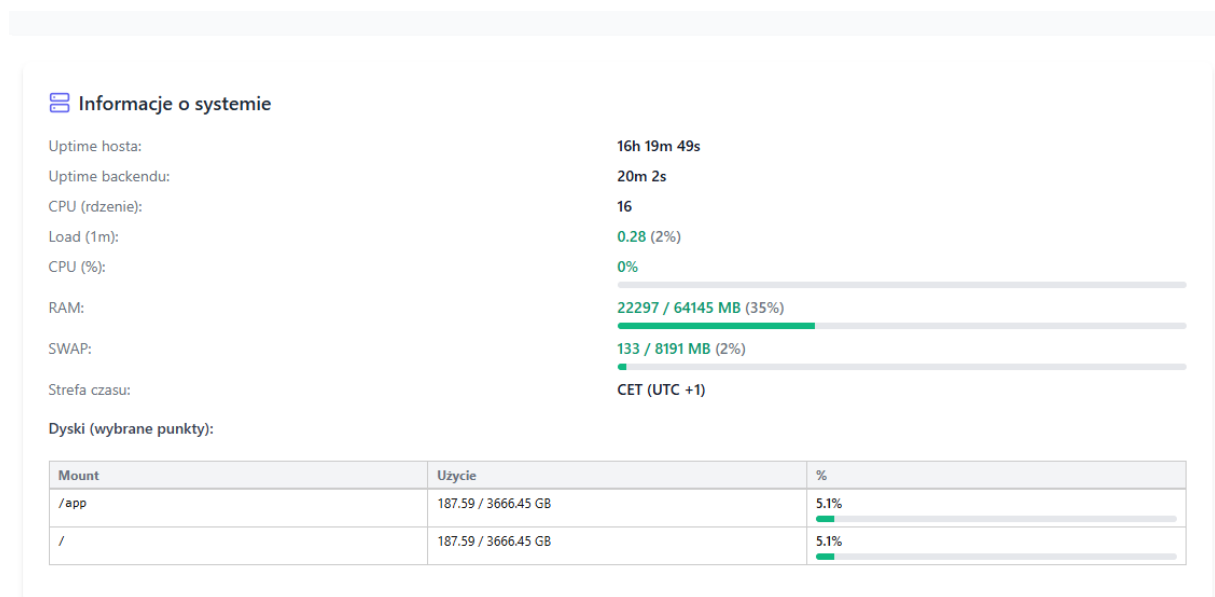
Wskaźniki w formie kół:



- **RAM, SWAP, GPU VRAM** (jeśli dostępne),
- **Dyski** – top 3 woluminów według rozmiaru.

Kolory zgodne z progami – szybki wgląd, co zbliża się do limitu.

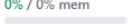
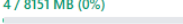
Informacje o systemie



- **Load (1m)** – średnie obciążenie z minuty (patrz słownik poniżej),
- **CPU (%)** – bieżące wykorzystanie,
- **RAM i SWAP** – zajętość oraz pasek postępu,
- **Dyski** – tabela z procentowym użyciem i paskiem postępu przy każdym montowaniu,
- **Strefa czasu.**

GPU

GPU

Idx	Model	Użycie	VRAM	Temp	Power	Procesy
0	NVIDIA GeForce RTX 5060 drv 575.57.08 GPU-0ad23a87-0444-d9b1-aade-300368c564ce	0% / 0% mem 	4 / 8151 MB (0%) 	42°C	8W / 145W	brak

Jeśli serwer ma GPU, zobaczysz tabelę:

- **Idx, Model** (i wersję sterownika),
- **Użycie GPU (%) i VRAM (MB i %)**,
- **Temperatura (°C)**,
- **Pobór mocy / limit (W)**,
- **Procesy** używające GPU (PID, nazwa, zajęcie VRAM).

Integracje

Tabela sond łączności:

Połączenia (sondy)

Nazwa	Host:Port	Schemat	Status	Latency	Błąd
qdrant	qdrant:6333	http	OK	1 ms	-
ollama	ollama:11434	http	OK	1 ms	-

- **Nazwa, Host:Port, Schemat, Status (OK/FAIL)**,
- **Latencja (ms)** – pomocna przy diagnozie opóźnień,
- **Błąd** – jeśli wystąpił.

4 Kolory i progi – jak interpretować

- **Zielony (OK)** – normalna praca.
- **Bursztynowy (WARN)** – zbliżasz się do zakresów nieoptymalnych:
 - CPU/RAM/GPU util/VRAM $\geq 85\%$,
 - SWAP $\geq 50\%$,
 - Dysk $\geq 80\%$,

- Temperatura GPU $\geq 80^{\circ}\text{C}$,
- Latencja ≥ 200 ms.
- **Czerwony (CRIT)** – wartości wysokie i potencjalnie groźne:
 - CPU/RAM/GPU util/VRAM $\geq 95\%$,
 - SWAP $\geq 80\%$,
 - Dysk $\geq 90\%$,
 - Temperatura GPU $\geq 90^{\circ}\text{C}$,
 - Latencja ≥ 500 ms.

Co robić?

- **CPU/RAM** stale wysoko: ogranicz równoległe zadania, rozważ skalowanie/zmianę priorytetów.
- **SWAP** wysoko: zamknij pamięciożerne procesy, to sygnał niedoboru RAM.
- **Dyski** > 90%: posprzątaj stare dane/logi lub powiększ wolumen.
- **GPU** wysoko lub gorąco: wstrzymaj kolejkę zadań, popraw chłodzenie, ogranicz batch-size.
- **Latencja** rośnie: sprawdź sieć, obciążenie usługi docelowej, ewentualne restarty/piki ruchu.

5 Częste scenariusze i szybkie rozwiązania

- **Ollama/Qdrant na czerwono**
Sprawdź „Szczegóły” → „Błąd”. Często chodzi o nieosiągalny endpoint lub restart usługi.
- **Brak sekcji GPU**
Serwer nie ma GPU lub aplikacja nie ma uprawnień/sterowników. Zgłoś do administratora.
- **Wartości skaczą**
Włącz **auto-odświeżanie 10 s** i obserwuj trend. Skoki pojedyncze są normalne; istotna jest **tendencja**.
- **Duża latencja**
Sprawdź, czy problem dotyczy jednej integracji czy wielu. Jedna – problem po stronie usługi; wiele – sieć lub host.
- **Chcę zgłosić incydent**
Kliknij **Eksport CSV** i dołącz plik do zgłoszenia (zawiera historię trendów z Twojej przeglądarki).

6 Słownik pojęć

- **CPU (%)** – ile mocy obliczeniowej procesora aktualnie używa system.
- **Load (1m)** – średnie obciążenie z ostatniej minuty. Dla 8-rdzeniowego CPU **8.0 $\approx$ 100%**.
- **RAM** – pamięć operacyjna; gdy jest prawie pełna, system zaczyna **swapować** (zwalnia).
- **SWAP** – „awaryjna” przestrzeń na dysku używana jak dodatkowa pamięć; jej wysoka używalność spowalnia system.

- **Dysk (%)** – zajętość woluminu; powyżej 90% ryzyko braku miejsca/logów.
- **GPU util (%)** – jak intensywnie pracuje karta graficzna (rdzenie).
- **VRAM** – pamięć karty graficznej; zbyt mała wolna VRAM ogranicza zadania AI.
- **Temperatura (°C)** – temperatura GPU; wysokie wartości grożą throttlingiem.
- **Power (W)** – pobór mocy GPU; zbliżanie się do limitu może ograniczać wydajność.
- **Latencja (ms)** – opóźnienie odpowiedzi usługi (niższe = lepiej).
- **Alive** – sygnał, że usługa odpowiada na proste zapytania kontrolne.
- **Endpoint** – adres i port usługi.

7 Prywatność i bezpieczeństwo

- Trendy (historia wykresów) są **zapisywane lokalnie** w Twojej przeglądarce.
- Żadne dodatkowe dane nie są wysyłane na zewnątrz.
- Surowe dane w „Szczegółach” pochodzą z backendu i nie zawierają sekretów.

8 Najlepsze praktyki

- Na co dzień ustaw **auto-odświeżanie 30 s**. W trakcie incydentu – **10 s**.
- Gdy chcesz „od zera” zobaczyć trendy – kliknij **Reset historii**.
- Jeśli coś wygląda źle – najpierw sprawdź **Baner alertów** i **Status usług**, potem przyjrzyj się **Trendom**.

9 FAQ (krótkie odpowiedzi)

Czy mogę zmienić progi kolorów?

Nie z poziomu strony. Progi są stałe – ewentualną zmianę koordynuje administrator.

Dlaczego wykresy „znikają” po restarcie przeglądarki?

Historię trzymamy lokalnie. Usunięcie danych przeglądarki lub inna przeglądarka = brak historii.

Nie widzę żadnych danych.

Sprawdź uprawnienie „**Podgląd Stanu Systemu**” i dostępność backendu. Użyj „Odśwież teraz”.

10 Kontakt i eskalacja

W zgłoszeniu podaj:

1. Czas problemu (z paska **Ostatnia aktualizacja**).
2. Zrzut ekranu sekcji **Status usług** i/lub **Baner alertów**.
3. **CSV** z trendami (przycisk **Eksport CSV**).
4. Krótki opis czynności, które wykonywałeś.

10. Pomoc i wsparcie

Sekcja pomocy i wsparcia została stworzona, aby dostarczyć użytkownikom i administratorom niezbędne narzędzia i informacje do rozwiązywania problemów oraz optymalnego wykorzystania platformy AsWiseAI.

Najczęściej zadawane pytania (FAQ)

Poniżej znajduje się lista najczęściej zadawanych pytań, która może pomóc w szybkim rozwiązaniu typowych problemów.

- **Jak dodać nowe dokumenty?**
 - Aby dodać dokumenty do bazy wiedzy, przejdź do strony Baza wiedzy w panelu nawigacyjnym. Możesz tam przeciągnąć i upuścić pliki lub wybrać je z dysku. Akceptowane formaty to PDF, PNG, JPG i TXT.
- **Dlaczego AI nie może odpowiedzieć na moje pytanie?**
 - Prawdopodobnie odpowiedź nie znajduje się w dokumentach, które zostały załadowane do bazy wiedzy. Model AI jest ograniczony do korzystania wyłącznie z dostarczonego kontekstu i nie ma dostępu do wiedzy zewnętrznej.
- **Czy moje dane są bezpieczne?**
 - Tak, AsWiseAI działa w modelu on-premise, co oznacza, że wszystkie Twoje dane (dokumenty, historia zapytań, wektory) są przechowywane lokalnie na Twojej infrastrukturze. Aplikacja nie wysyła żadnych danych do zewnętrznych usług chmurowych.
- **Czym różni się "AI Fakt" od "AI Agent Analityczny"?**
 - **AI Fakt** to tryb do zadawania konkretnych, pojedynczych pytań. Odpowiedzi są zwięzłe, a każda informacja jest poparta cytatem ze źródła.
 - **AI Agent Analityczny** jest przeznaczony do prowadzenia dłuższych, płynnych konwersacji, gdzie AI pamięta kontekst rozmowy. Odpowiedzi również zawierają cytaty i linki do źródeł.
- **Jak zresetować hasło, jeśli go zapomniałem?**
 - Przejdź na stronę logowania i wybierz opcję Nie pamiętasz hasła?. Podaj swoją nazwę użytkownika, a na Twój adres e-mail zostanie wysłany link do zresetowania hasła.

Kontakt z pomocą techniczną

W przypadku problemów, których nie da się rozwiązać za pomocą tej dokumentacji, skontaktuj się z pomocą techniczną.

Adres e-mail: biuro@entersoft.pl

Słowniczek pojęć

- **LLM (Large Language Model):** Duży model językowy, np. Mistral lub Llama, który przetwarza i generuje naturalny język. W AsWiseAI są one uruchamiane lokalnie za pomocą usługi Ollama.
- **On-premise:** Model wdrożenia, w którym oprogramowanie jest instalowane i działa na serwerach klienta, a nie w zewnętrznej chmurze. Zapewnia to pełną kontrolę nad danymi.
- **RAG (Retrieval-Augmented Generation):** Technika, którą stosuje AsWiseAI. Model AI najpierw przeszukuje bazę wiedzy (Twoje dokumenty), a następnie używa znalezionych fragmentów jako kontekstu do wygenerowania odpowiedzi.
- **Embedding:** Proces przekształcania tekstu w wektor liczbowy. Jest to kluczowy element, który pozwala na semantyczne przeszukiwanie dokumentów w bazie wektorowej.
- **OCR (Optical Character Recognition):** Technologia używana do rozpoznawania tekstu na obrazach (np. zeskanowanych dokumentach).
- **Qdrant:** Baza danych wektorowych, która przechowuje embeddingi. Jest to kluczowy komponent AsWiseAI, umożliwiający błyskawiczne i precyzyjne przeszukiwanie Twoich dokumentów w oparciu o ich znaczenie.
- **Vector Database:** Specjalistyczna baza danych, która przechowuje dane w postaci wektorów. Umożliwia ona wyszukiwanie informacji na podstawie podobieństwa semantycznego, a nie tylko słów kluczowych. W AsWiseAI tę rolę pełni Qdrant.
- **Chunk:** Pojedynczy fragment tekstu, na który dzielone są Twoje dokumenty. Dzielenie na mniejsze fragmenty (chunkowanie) jest niezbędne, aby model AI mógł efektywnie je przetwarzać i generować spójne odpowiedzi.
- **Prompt:** Tekstowe zapytanie lub instrukcja, którą przekazujesz modelowi AI, aby uzyskać konkretną odpowiedź lub wykonać zadanie. W AsWiseAI jest to Twoje pytanie, które kierujesz do swoich dokumentów.
- **Prompt Systemowy:** Jest to zestaw instrukcji, które definiują zachowanie modelu AI, np. jego rolę, zasady bezpieczeństwa, czy styl odpowiedzi. Administratorzy mogą tworzyć i edytować różne wersje tych promptów.
- **Wersjonowanie promptów:** Funkcjonalność, która pozwala na tworzenie wielu wersji promptów systemowych i bezpieczne ich testowanie w Piaskownicy (Sandbox).
- **Few-shot examples:** Przykłady par pytań i odpowiedzi, które służą do uczenia modelu AI, w jaki sposób ma formatować i prezentować dane w odpowiedziach.
- **Celery:** Kolejka zadań, która w AsWiseAI odpowiada za wykonywanie długotrwałych operacji w tle, takich jak indeksowanie dużych plików. Dzięki temu interfejs użytkownika pozostaje płynny i responsywny.

- **Tokeny uwierzytelniające:** Małe fragmenty danych, przechowywane w pamięci lokalnej przeglądarki, które umożliwiają utrzymanie sesji użytkownika. AsWiseAI używa ich zamiast inwazyjnych plików cookie.
- **Historia zapytań (Historia QA):** Moduł, w którym przechowywane są wszystkie zapytania i odpowiedzi z interakcji z AI. Dane te są przechowywane lokalnie.
- **Powiadomienia:** System informowania użytkownika o statusie długotrwałych operacji, np. o pomyślnym przetworzeniu pliku lub o błędach. Są one widoczne po kliknięciu ikony dzwonka.
- **Tagi:** Opcjonalne etykiety, które można dodać do dokumentów podczas ich przesyłania. Pomagają one w późniejszym filtrowaniu i kategoryzacji plików.
- **Hashowanie haseł:** Twoje hasła są przechowywane w formie haszowanej, co uniemożliwia ich odczytanie i zwiększa bezpieczeństwo.
- **RBAC (Role-Based Access Control):** Kontrola dostępu oparta na rolach, która zarządza dostępem do poszczególnych funkcji aplikacji za pomocą ról i uprawnień. Dzięki temu tylko upoważnieni użytkownicy mają dostęp do wrażliwych danych i funkcji.
- **Responsywność:** Aplikacja działa płynnie na różnych urządzeniach, zarówno na komputerach stacjonarnych, jak i na urządzeniach mobilnych, takich jak smartfony i tablety. Interfejs użytkownika automatycznie dostosowuje swój układ do mniejszych ekranów, a nawigacja jest zoptymalizowana pod kątem dotyku.
- **Skalowalność:** Zdolność systemu AsWiseAI do efektywnego działania i obsługi rosnącego obciążenia, takiego jak zwiększona liczba użytkowników, większa ilość przetwarzanych dokumentów czy bardziej złożone zapytania. System został zaprojektowany tak, aby można było łatwo zwiększać jego moc obliczeniową w miarę potrzeb.
- **Hash (skrót kryptograficzny):** Unikalny, stałej długości ciąg znaków, generowany na podstawie dowolnie dużego pliku lub ciągu danych. W AsWiseAI hashowanie jest wykorzystywane do weryfikacji integralności plików oraz do bezpiecznego przechowywania haseł, ponieważ oryginalnych danych nie da się odzyskać z hasha.
- **Plik CSV (Comma-Separated Values):** Plik tekstowy, w którym dane są przechowywane w formie tabelarycznej, a poszczególne wartości oddzielone są separatorami/przecinkami. W AsWiseAI format ten jest wykorzystywany do importowania i eksportowania np. przykładów "few-shot" w panelu zarządzania promptami.

11. Rozwiązywanie problemów i dobre praktyki

Ta sekcja zawiera odpowiedzi na najczęstsze problemy techniczne oraz zbiór dobrych praktyk, które pomogą Ci zoptymalizować działanie i niezawodność systemu AsWiseAI, ze szczególnym uwzględnieniem modułu Automatyzacji.

Najczęstsze problemy i rozwiązania (FAQ Techniczne)

Problem: Mój harmonogram automatyzacji nie uruchamia się o czasie.

- **Rozwiązanie w 3 krokach:**

1. **Sprawdź status zadania w UI:** Upewnij się, że w panelu **Automatyzacja**, przy zadaniu, które Cię interesuje, przełącznik statusu jest w pozycji aktywnej (zielony).
2. **Zweryfikuj logi celery-beat:** Poproś administratora o sprawdzenie logów kontenera celery-beat (docker compose logs -f celery-beat). Powinny się w nich cyklicznie pojawiać komunikaty o wysyłaniu zadań, np. Scheduler: Sending due task.... Jeśli logi są puste lub zawierają błędy, usługa harmonogramu może wymagać restartu.
3. **Sprawdź Dziennik Zdarzeń:** Kliknij ikonę historii (🕒) przy zadaniu. Jeśli ostatnie wpisy mają status SKIPPED_LOCKED, oznacza to, że poprzedni cykl przetwarzania wciąż trwa, a system – zgodnie z polityką współbieżności SKIP – inteligentnie pomija kolejne uruchomienie, aby uniknąć przeciążenia.

Problem: Pliki w monitorowanym katalogu nie są przetwarzane, mimo że harmonogram się uruchamia.

- **Rozwiązanie w 4 krokach:**

1. **Sprawdź logi celery-worker:** To najważniejszy krok. Logi kontenera celery-worker (docker compose logs -f celery-worker) zawierają szczegółowe informacje o każdym przetwarzanym pliku. Znajdziesz tam komunikaty o sukcesie (succeeded) lub dokładny opis błędu (Traceback), który wystąpił.
2. **Sprawdź wzorzec nazwy pliku:** Jeśli w konfiguracji zadania użyłeś "Wzorca nazwy pliku" (np. RAPORT_*.pdf), upewnij się, że nazwy plików w katalogu dokładnie pasują do tego wzorca.
3. **Sprawdź, czy plik nie jest duplikatem:** System identyfikuje duplikaty na podstawie treści pliku (hash), a nie jego nazwy. Jeśli wrzuciłeś plik, który już kiedyś został przetworzony, system go pominie (zgodnie z ustawieniem "Nadpisz istniejące pliki"). W Dzienniku Zdarzeń zobaczysz wtedy wpis Pominęto duplikatów: 1. To nie błąd, a poprawne działanie systemu.
4. **Sprawdź uprawnienia do plików:** Upewnij się, że system (użytkownik, w kontekście którego działa Docker) ma uprawnienia do odczytu i zapisu w katalogu auto_ingest oraz jego podkatalogach na serwerze.

Problem: Edytowałem istniejące zadanie w panelu, ale ono wciąż działa według starych zasad.

- **Rozwiązanie:** Harmonogram (celery-beat) wczytuje konfigurację wszystkich zadań z bazy danych **tylko raz, podczas swojego startu**. Aby nowo zapisane zmiany (np. zmiana interwału,

ścieżki czy tagów) zostały uwzględnione, usługa celery-beat musi zostać zrestartowana. Można to zrobić za pomocą polecenia: `docker compose restart celery-beat`.

Problem: Logowanie przez Active Directory nie działa.

- **Rozwiązanie:**
 1. **Użyj przycisku "Testuj Połączenie"** w panelu Integracji. Pokaże on szczegółowy komunikat o ewentualnym błędzie.
 2. **Sprawdź zaporę sieciową (firewall):** Upewnij się, że serwer AsWiseAI ma możliwość komunikacji z Twoim kontrolerem domeny na porcie LDAPS (domyślnie 636).
 3. **Zweryfikuj dane konta serwisowego:** Sprawdź dokładnie poprawność wpisanej nazwy Bind DN oraz hasła dla konta technicznego, którego AsWiseAI używa do połączenia.
 4. **Sprawdź ścieżkę Search Base DN:** Upewnij się, że wskazuje ona na jednostkę organizacyjną (OU) w AD, w której znajdują się konta użytkowników, którzy mają mieć dostęp do systemu.

Dobre praktyki i optymalizacja (Pro-Tips)

- **Zaczynaj od małych partii i krótkich interwałów:** Konfigurując nowe zadanie automatyzacji, ustaw "Uruchamiaj co (minuty)" na 1, a "Ilość plików na cykl" na niską wartość (np. 5). Pozwoli Ci to szybko i bezpiecznie przetestować całą logikę. Gdy potwierdzisz, że wszystko działa zgodnie z oczekiwaniami, możesz zwiększyć interwał do bardziej rozsądnej wartości (np. 15 minut) i podnieść limit plików.
- **Stosuj precyzyjne wzorce nazw plików:** Zamiast ogólnego *.pdf, używaj bardziej szczegółowych wzorców, jak FAKTURA_*.pdf lub raport-*-2025.xlsx. Zapobiegnie to przypadkowemu przetworzeniu plików tymczasowych (np. ~WRD000.tmp) lub niepowiązanych dokumentów, które mogą pojawić się w monitorowanym katalogu.
- **Mądrze zarządzaj akcjami "Po sukcesie":**
 - Przenieś do 'processed': Jest to najbezpieczniejsza i zalecana opcja. Masz wgląd w to, które pliki zostały pomyślnie przetworzone.
 - Usuń plik: Używaj tej opcji z ostrożnością, upewniając się, że masz kopię zapasową plików źródłowych.
 - Zostaw: Przydatne do testów lub w sytuacji, gdy inny system jest odpowiedzialny za archiwizację plików po ich przetworzeniu przez AsWiseAI.
- **Stwórz spójny standard tagowania:** Ustal w ramach organizacji schemat dla "Domyślnych Tagów JSON". Dobre, spójne tagi są bezcenne dla późniejszej analizy przez AI Agent. Na przykład, wszystkie faktury mogą mieć tag ["faktura", "finanse"], a umowy ["umowa", "prawne"]. Pozwoli to agentom na precyzyjniejsze zawężanie kontekstu.

Praktyczny przykład: Automatyzacja faktur kosztowych

1. **Cel:** Chcemy, aby wszystkie skany faktur kosztowych były automatycznie przetwarzane i tagowane rokiem ich otrzymania.

2. **Konfiguracja folderu:** Na serwerze, w katalogu data_dev/auto_ingest/, tworzymy strukturę, np. Faktury/Koszty/.
3. **Konfiguracja zadania w UI:**
 - **Ścieżka:** Faktury/Koszty
 - **Skanuj rekursywnie:** Włączone (aby system mógł obsługiwać podkatalogi, np. .../Koszty/2025/Styczen/).
 - **Wzorzec nazwy pliku:** FK_*.pdf (dla plików typu "Faktura Kosztowa").
 - **Domyślne Tagi:** ["faktura", "kosztowa"]
 - **Nadpisz istniejące pliki:** Włączone (na wypadek, gdyby ktoś wgrał poprawioną wersję skanu tej samej faktury).
 - **Email dla błędów:** ksiegowosc@twojafirma.pl
4. **Działanie:** Teraz wystarczy, że uprawniona osoba wrzuci plik o nazwie np. FK_01-01-2025_dostawca.pdf do folderu /mnt/AsWiseAI/data_dev/auto_ingest/Faktury/Koszty/. Automat go wykryje, przetworzy i doda do bazy wiedzy z tagami ["faktura", "kosztowa"], gotowy do analizy.
5. **Interakcja z AI:** Po przetworzeniu, możesz zadać pytanie w trybie AI Agent, np.: "Wylistuj wszystkie faktury kosztowe od 'dostawca' i podsumuj ich kwoty netto.".

12. Zaawansowane Funkcje Interfejsu Czatów

1. Ręczny Wybór Agent Specjalistycznego

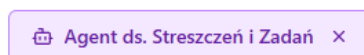


Obok pola do wprowadzania tekstu znajduje się przycisk z ikoną Bota, który pozwala na ręczne wybranie jednego z dostępnych Agentów Specjalistycznych.

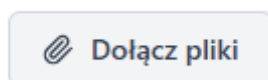
Jak działa?: Wybór agenta z tej listy (np. "Agent HR") jest **manualnym nadpisaniem Orkiestratora**. System nie będzie już analizował pytania, aby dobrać specjalistę, ale od razu przekaże Twoje polecenie do wybranego przez Ciebie Agent. Jest to tzw. **Tryb Wyzwalany**.

Korzyść: Daje Ci to pełną kontrolę i pozwala świadomie korzystać z wiedzy konkretnego, wirtualnego eksperta, gdy wiesz dokładnie, jakiego rodzaju analizy potrzebujesz.

Wskazówka wizualna: Po wybraniu Agent, pole tekstowe zmieni obramowanie (np. na fioletowe), a jego nazwa zostanie wyświetlona, informując Cię, że działasz w trybie specjalistycznym.



2. Dołączanie Plików do Zapytania (Kontekst Tymczasowy)



Ikona spinacza pozwala na dołączenie jednego lub więcej plików (PDF, obrazów, plików tekstowych) bezpośrednio do Twojego zapytania.

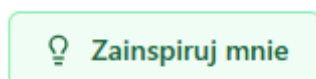
Jak działa?: System odczyta treść załączonych plików (w razie potrzeby używając technologii OCR) i potraktuje ją jako **tymczasowy, jednorazowy kontekst** tylko dla tego jednego, konkretnego pytania.

Kluczowa różnica: Pliki dołączone w ten sposób **NIE SĄ** dodawane do trwałej Bazy Wiedzy. Służą one wyłącznie do analizy "w locie".

Przykład użycia: Możesz załączyć plik PDF z nową umową i od razu zapytać: *"Streść klauzule dotyczące kar umownych w tym dokumencie"*, bez potrzeby wcześniejszego, trwałego indeksowania pliku w systemie.

4. Sugestie AI ("Zainspiruj mnie")

Przycisk z ikoną żarówki służy do generowania inteligentnych sugestii pytań.



Jak działa?: Po kliknięciu, AsWiseAI losowo analizuje fragmenty Twojej bazy wiedzy i na ich podstawie generuje 3-4 przykładowe, otwarte pytania, które mogą Cię zainspirować do dalszej eksploracji danych.

Korzyść: Jest to doskonałe narzędzie, gdy nie jesteś pewien, od czego zacząć lub chcesz odkryć, jakie interesujące informacje mogą kryć się w Twoich dokumentach.

13. Transkrypcja Plików Audio i Wideo

AsWiseAI został wyposażony w zaawansowany moduł do automatycznej transkrypcji, który przekształca mowę z plików audio i wideo na tekst. Dzięki integracji z technologią OpenAI Whisper, system jest w stanie z dużą precyzją rozpoznawać mowę w wielu językach, tworząc tekst, który może być następnie przeszukiwany lub dodany do bazy wiedzy. Moduł ten jest dostępny w menu głównym pod pozycją **Transkrypcje**.

Transkrypcja Plików Audio/Wideo

Prześlij pliki, wybierz tryb przetwarzania i opcjonalnie zaplanuj zadanie na później.

Pliki do przetworzenia


Wybierz pliki lub przeciągnij i upuść
MP4, MOV, MP3, WAV, etc.

Tryb Przetwarzania

☐  **Ekspres (GPU)**
Szybciej, wymaga zasobów GPU

☒  **Standard (CPU)**
Wolniej, mniejsze obciążenie

Moment Uruchomienia

☒  **Natychmiast**
Zadanie trafi do kolejki od razu

☐  **Zaplanuj**
Uruchom o wskazanej godzinie

Ustawienia Zaawansowane

Model Whisper 

Język 

Temperatura: 0 

Medium (zalecany) 

pl



Opcje Dekodowania

Beam Size 

Patience 

Ignorowane tokeny 

5

brak

1

Wstępny prompt (kontekst) 

Podaj listę nazwisk, akronimów...

Opcje Przetwarzania i Wyjścia

☒ Wytnij fragmenty ciszy (VAD) ☐ Znaczniki czasu dla słów ☐ Identyfikuj mówców (Diarize)

☒ Dodaj wynik do bazy wiedzy AsWiseAI
Odmarczenie tej opcji spowoduje, że transkrypt będzie dostępny tylko do pobrania i nie będzie przeszukiwany przez AI.

Rozpocznij Transkrypcję

Historia Zadań Transkrypcji

PLIK	STATUS	TRYB	ZŁOŻONO	UROŚNIO	AKCJE
------	--------	------	---------	---------	-------

1. Zlecenie Nowej Transkrypcji

Główny panel modułu pozwala na przesłanie plików i precyzyjne skonfigurowanie zadania transkrypcji. Proces składa się z kilku kroków:


EnterSoft

Firma Informatyczna EnterSoft
e-mail: biuro@entersoft.pl, entersoft.pl, kom.: 696 481 842

Str.: 90/99

1.1. Wybór Plików Możesz wybrać jeden lub więcej plików audio/wideo, przeciągając je na wyznaczony obszar lub klikając w celu otwarcia okna dialogowego. System obsługuje większość popularnych formatów (np. MP4, MOV, MP3, WAV).

1.2. Tryb Przetwarzania Masz do wyboru dwie opcje, które wpływają na szybkość i obciążenie serwera:

- **Ekspres (GPU):** Rekomendowany tryb, który wykorzystuje moc obliczeniową karty graficznej (GPU) do znacznie szybszego przetwarzania. Należy go używać, jeśli serwer jest wyposażony w odpowiednią kartę NVIDIA.
- **Standard (CPU):** Tryb, który używa wyłącznie procesora (CPU). Jest wolniejszy, ale nie wymaga użycie zasobów pamięci GPU.

1.3. Moment Uruchomienia

- **Natychmiast:** Zadanie zostanie dodane do kolejki do wykonania tak szybko, jak to możliwe.
- **Zaplanuj:** Pozwala na odroczenie zadania i uruchomienie go o konkretnej, wybranej dacie i godzinie. Jest to przydatne do planowania zasobochłonnych zadań na godziny nocne.

1.4. Opcje Zaawansowane Sekcja ta pozwala na precyzyjne dostrojenie parametrów transkrypcji w celu uzyskania jak najlepszej jakości:

- **Model Whisper:** Wybór wielkości modelu AI. Większe modele (medium, large) są dokładniejsze, ale działają wolniej i wymagają więcej zasobów pamięci. Mniejsze (tiny, base, small) są szybsze, ale mogą popełniać więcej błędów.
- **Język:** Możesz ręcznie wpisać kod języka (np. pl dla polskiego, en dla angielskiego itd.), aby zwiększyć dokładność. Jeśli pole pozostanie puste, model spróbuje automatycznie wykryć język.
- **Wstępny prompt (kontekst):** To bardzo użyteczna funkcja. Możesz tu wpisać listę specyficznych słów, nazwisk, akronimów lub żargonu technicznego, które występują w nagraniu. Pomoże to modelowi AI poprawnie je rozpoznać.
- **Inne opcje:** Możesz również dostosować zaawansowane parametry, takie jak Temperatura (kreatywność modelu), Beam Size (wpływa na spójność tekstu) czy Wytnij fragmenty ciszy (VAD) (znacząco przyspiesza proces przez usuwanie ciszy z nagrania).

1.5. Dodaj wynik do bazy wiedzy To kluczowa opcja decydująca o przeznaczeniu wyniku:

- **Opcja ZAZNACZONA (domyślnie):** Wygenerowany tekst zostanie automatycznie podzielony na fragmenty, zwektoryzowany i dodany do bazy wiedzy Qdrant. Będzie on przeszukiwalny przez AI Fakt i AI Agent. Nazwa oryginalnego pliku zostanie poprzedzona prefiksem [Transkrypt].
- **Opcja ODZNACZONA:** Wygenerowany tekst zostanie zapisany jako plik .txt na serwerze i **nie będzie** dodany do bazy wiedzy. Taki wynik będzie dostępny wyłącznie do pobrania z panelu historii zadań.

2. Historia Zadań Transkrypcji

Poniżej formularza znajduje się tabela z historią wszystkich zleconych zadań transkrypcji. Pozwala ona na bieżąco monitorować ich status i zarządzać wynikami.

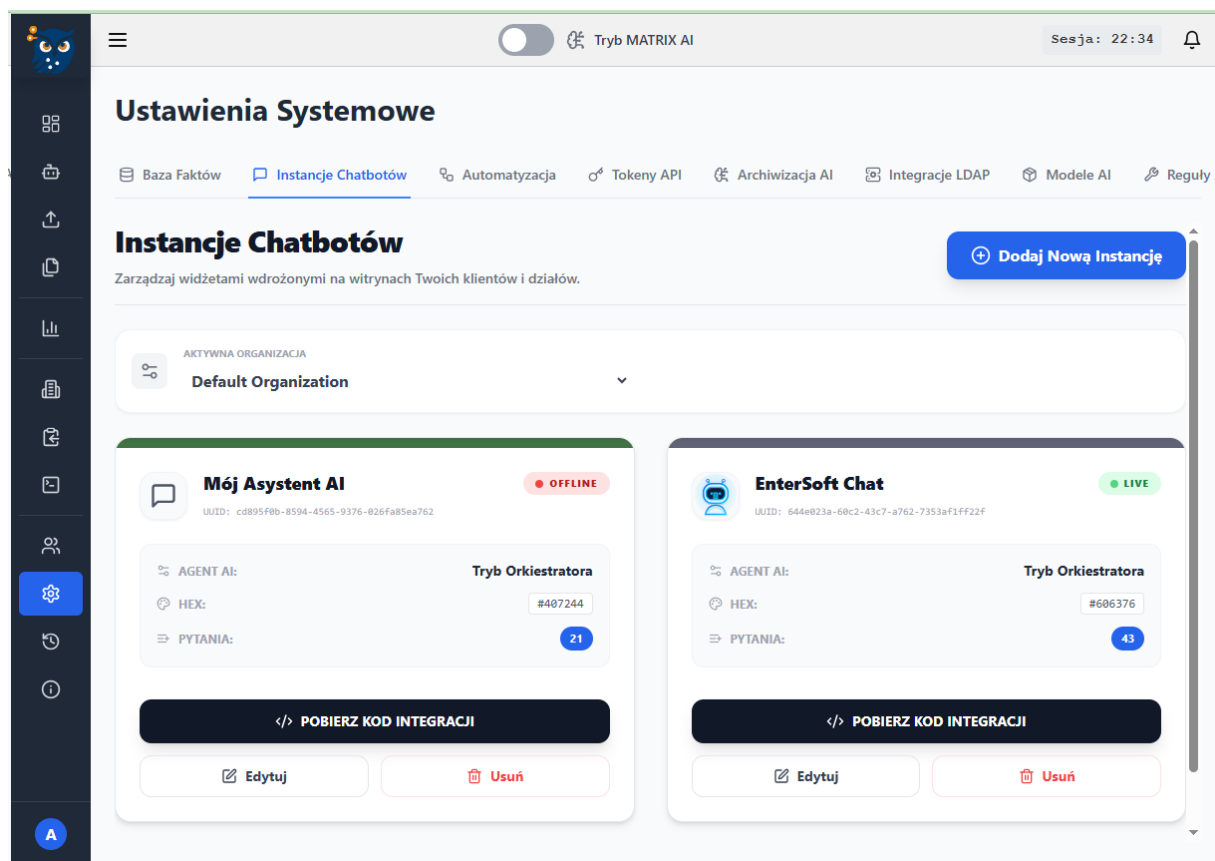
- **Statusy Zadań:**

- W kolejce: Zadanie czeka na wolny zasób (worker CPU/GPU).
 - Przetwarzanie: Zadanie jest w toku. Wyświetlany jest pasek postępu oraz informacja o bieżącym etapie (np. "Ekstrakcja audio...", "Transkrypcja w toku...").
 - Ukończono: Proces zakończył się sukcesem.
 - Błąd: Wystąpił problem podczas przetwarzania.
 - Anulowano: Zadanie zostało zatrzymane przez użytkownika.
- **Dostępne Akcje:**
 - **Pobierz:** Pojawia się dla ukończonych zadań, które **nie były** dodawane do bazy wiedzy. Umożliwia pobranie pliku .txt z transkryptem.
 - **Ponów:** Dostępna dla zadań, które zakończyły się błędem. Zleca zadanie do wykonania od nowa.
 - **Zatrzymaj:** Pojawia się dla zadań w kolejce lub w trakcie przetwarzania. Umożliwia natychmiastowe anulowanie zadania.
 - **Usuń:** Dostępna dla zadań zakończonych (sukcesem, błędem lub anulowanych). Usuwa wpis z historii oraz wszystkie powiązane pliki z serwera (oryginalne nagranie, plik wynikowy).

14. Zarządzanie Instancjami Chatbotów (Widżety WWW)

Panel **Instancje Chatbotów** to zaawansowane narzędzie pozwalające na "wystawienie" wiedzy zgromadzonej w AsWiseAI na zewnątrz – w formie interaktywnego widżetu na Twojej stronie internetowej, w intranecie lub w portalu dla klientów.

Dzięki temu modułowi możesz stworzyć nieskończoną liczbę niezależnych asystentów. Każdy z nich może mieć inny wygląd, inną bazę pytań startowych oraz odpowiadać w innym tonie, będąc jednocześnie w pełni zintegrowanym z Twoją wewnętrzną bazą wiedzy.



1. Tworzenie i Edycja Chatbota

Aby stworzyć nowego bota, przejdź do zakładki **Ustawienia Systemowe -> Instancje Chatbotów** i kliknij przycisk **Dodaj Nową Instancję**.

Otworzy się panel konfiguracyjny podzielony na cztery logiczne sekcje.

Edycja Konfiguracji Chatbota
×

WYGLĄD I TOŻSAMOŚĆ

Nazwa wyświetlana w nagłówku

EnterSoft Chat

Kolor przewodni (Branding)

#586376

LOGO (NAGŁÓWEK)

Zmień

OBRAZ SUGESTII

Zmień

AWATAR AI

Zmień

Wiadomość powitalna na start

Cześć! W czym mogę Ci dziś pomóc?
...?

INTELEGENCJA I ZACHOWANIE

☒
Losuj pytania startowe
Użytkownik zobaczy losowy zestaw z Twojej bazy przy każdym odliwieżeniu.

Własne instrukcje systemowe

Np.: Odpowiadaj jako ekspert od ubezpieczeń. Bądź bardzo precyzyjny i unikaj potocznego słownictwa...

Instrukcje te są dodawane do System Promptu agenta przed każdym zapytaniem.

☐
BAZA SUGESTII

Import

Eksport

Generuj z AI

Dodaj

☐
1. Jakże są zastosowania analizy w trakcie wykonania harmonogramu

☐
2. Czy możliwe jest określenie indywidualnych ustawień dla param

☐
3. W jaki sposób zmieniam wartości parametru w kamerze według

☐
4. Jak technologie harmonogramowania mogą pomóc w organiza

KONFIGURACJA SILNIKA I DOSTĘPU

Dedykowany Agent (Mózg AI)

-- Automatyczny (Orkiestrator) --

Token Serwisowy (Klucz Analityki)

NEW

Status Widoczności Widżetu

AKTYWNY (Publiczny na stronie)

Dozwolone domeny (CORS)

https://entersoft.pl

Anuluj

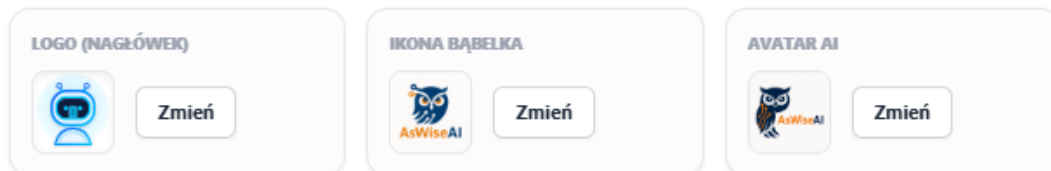
Zatwierdź i Zapisz Konfigurację

Sekcja A: Wygląd i Tożsamość

W tej sekcji decydujesz, jak bot będzie prezentował się użytkownikom na stronie:

- Nazwa wyświetlana:** Nagłówek okna czatu (np. "Asystent Sprzedaży", "Pomoc Techniczna").

- **Kolor przewodni:** Główny kolor interfejsu (w formacie HEX), który możesz dopasować do identyfikacji wizualnej (brandbooka) swojej firmy.
- **Personalizacja Graficzna:** Możesz wgrać trzy niezależne grafiki:
 - **Logo (Nagłówek):** Ikonka widoczna na górnym pasku otwartego czatu.
 - **Ikona bąbelka:** Grafika wyświetlana w rogu ekranu, zachęcająca do otwarcia czatu.
 - **Avatar AI:** Zdjęcie lub ikona, która pojawia się przy każdej odpowiedzi wygenerowanej przez bota.



- **Wiadomość powitalna na start:** Pierwsza wiadomość, którą użytkownik zobaczy natychmiast po otwarciu okna czatu.

Sekcja B: Inteligencja i Zachowanie

Tu nadajesz swojemu botowi unikalny charakter:

- **Losuj pytania startowe:** Po włączeniu tej opcji, system będzie losowo dobierał zestaw pytań z Twojej Bazy Sugestii przy każdym odświeżeniu strony przez użytkownika.
- **Własne instrukcje systemowe:** Ukryte polecenia modyfikujące zachowanie AI. Możesz tu wpisać np. *"Odpowiadaj zawsze w sposób bardzo oficjalny, używając zwrotów grzecznościowych. Jesteś ekspertem z zakresu prawa pracy."* System doklei tę instrukcję do każdego zapytania jako prompt systemowy, który stanowi instrukcję dla LLM'a.

Sekcja C: Baza Sugestii

Aby ułatwić użytkownikom rozpoczęcie rozmowy, możesz zdefiniować gotowe pytania startowe, które pojawią się nad polem tekstowym czatu.

- **Ręczne dodawanie:** Użyj przycisku *Dodaj nowe*, aby wpisać pytanie ręcznie.
- **Import / Eksport:** Możesz masowo zarządzać pytaniami ładując je ze spreparowanego pliku .csv lub .txt.
- **Generuj z AI:** System może samodzielnie wymyślić propozycje pytań!

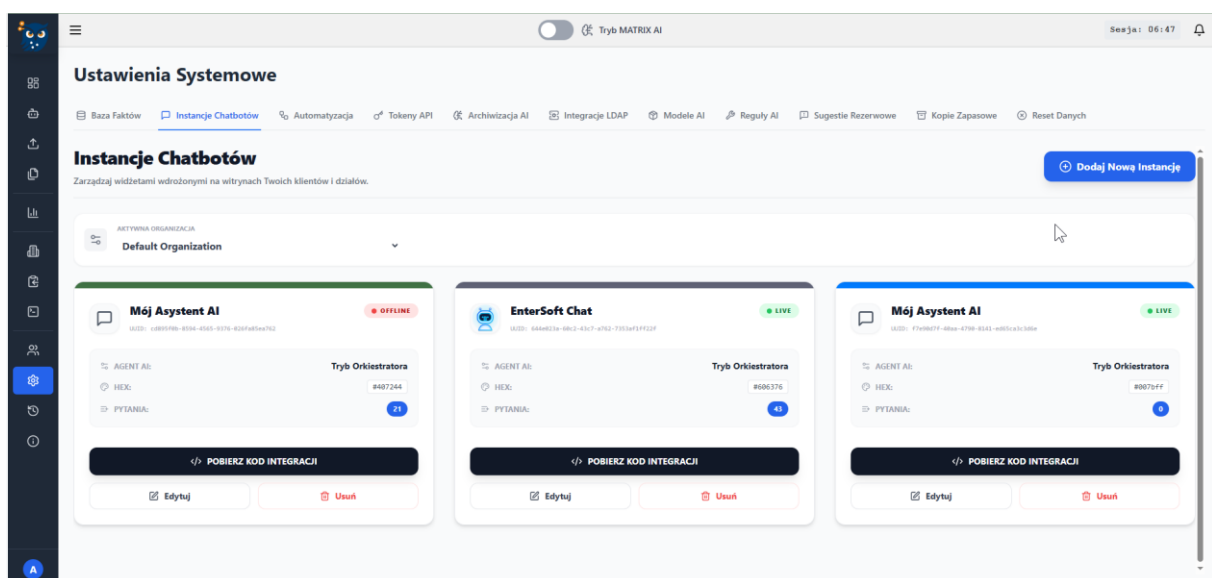
Wystarczy kliknąć przycisk **Generuj z AI** i podać, ile pytań chcesz uzyskać (od 1 do 50). System przeanalizuje dokumenty wgrane do Twojej Bazy Wiedzy i na ich podstawie zaproponuje najbardziej sensowne zagadnienia do rozmowy.

Sekcja D: Konfiguracja Silnika i Dostępu

Krytyczna sekcja odpowiadająca za techniczne działanie bota i bezpieczeństwo:

- **Dedykowany Agent (Mózg AI):** Możesz przypisać chatbota do konkretnego specjalisty z Twojej Macierzy Agentów (np. tylko do Agenta Prawnego). Pozostawienie opcji *Automatyczny (Orkiestrator)* sprawi, że bot będzie sam dobierał eksperta do zadanego pytania.
- **Token Serwisowy:** Wymagany klucz bezpieczeństwa. Każdy bot musi mieć przypisany aktywny token API, dzięki któremu jego operacje są autoryzowane i zapisywane w logach analitycznych (patrz TOKENY API).
- **Status Widoczności:** Szybki włącznik/wyłącznik (LIVE / OFFLINE). Wyłączenie bota natychmiastowo zablokuje możliwość komunikacji na wszystkich zewnętrznych stronach WWW.
- **Dozwolone domeny (CORS):** Zabezpieczenie przed nieautoryzowanym skopiowaniem Twojego bota na inne strony. Wpisz tu adresy (np. <https://twoja-strona.pl>), na których widget ma prawo działać.

2. Zarządzanie na Liście i Kopiowanie Kodu



Wszystkie utworzone boty widoczne są na przejrzystej liście w formie kafelków. Na każdym z nich widzisz błyskawiczne podsumowanie konfiguracji (przypisany agent, liczba pytań startowych, kolor) oraz obecny status aktywności (migająca zielona dioda **LIVE**).

Jak wdrożyć bota na stronę?

1. Na kafelku wybranego bota kliknij duży, czarny przycisk **POBIERZ KOD INTEGRACJI**.
2. System automatycznie wygeneruje i skopiuje do Twojego schowka gotowy fragment kodu HTML (<link> oraz <script>).
3. Przekaż ten kod osobie zarządzającej Twoją stroną internetową lub wklej go samodzielnie w sekcji <body> witryny. Widget natychmiast pojawi się na stronie.

Innowacyjne rozwiązania IT dla Twojego biznesu

Od dedykowanych systemów klasy ERP po zaawansowaną sztuczną inteligencję. Zwiększ efektywność swojej firmy z naszymi autorskimi rozwiązaniami.

Skontaktuj się z nami



Oprogramowanie Dedykowane

Tworzymy niezawodne systemy skrojone na miarę, idealnie dopasowane do specyfiki procesów logistycznych i produkcyjnych w Twojej firmie.



Nowość

Wdrożenia AI (AsWiseAI)

Automatyzacja, analiza dokumentów (RAG) i agenci konwersacyjni bazujący wyłącznie na wewnętrznej wiedzy Twojej organizacji. Pełne bezpieczeństwo On-Premise.



Utrzymanie i Wsparcie

Zapewniamy ciągłość działania Twoich systemów serwerowych oraz błyskawiczną reakcję naszego zespołu wsparcia technicznego i DevOps.



Innowacyjne rozwiązania IT dla Twojego biznesu

Od dedykowanych systemów klasy ERP po zaawansowaną inteligencję. Zwiększ efektywność swojej firmy z naszymi rozwiązaniami.

Skontaktuj się z nami



Oprogramowanie Dedykowane

Tworzymy niezawodne systemy skrojone na miarę, idealnie dopasowane do specyfiki procesów logistycznych i produkcyjnych w Twojej firmie.



Nowość

Wdrożenia AI (AsWiseAI)

Automatyzacja, analiza dokumentów (RAG) i agenci konwersacyjni bazujący wyłącznie na wewnętrznej wiedzy Twojej organizacji. Pełne bezpieczeństwo On-Premise.

EnterSoft Chat

Cześć! W czym mogę Ci dziś pomóc? ..?

Jak ważna jest funkcja "Resetuj datę ostatniego uruchomienia" w kontekście zarządzania harmonogramami?

Jak zdefiniowanie statusu transakcji wpływa na jakość wyników w raporcie?

Czy możliwe jest automatyczne zmienianie parametrów w kamerach w oparciu o harmonogram zadań?

Zadaj pytanie...

Oryginalną reakcję naszego zespołu wsparcia technicznego i DevOps.



15. Podsumowanie

Dziękujemy za zapoznanie się z dokumentacją platformy AsWiseAI. Mamy nadzieję, że zawarte w niej informacje ułatwią Ci rozpoczęcie pracy i pozwolą w pełni wykorzystać potencjał naszej platformy analitycznej.

Życzymy efektywnej i owocnej pracy z AsWiseAI – Twoim zaufanym asystentem do analizy dokumentów.

W przypadku jakichkolwiek pytań, wątpliwości lub sugestii, prosimy o kontakt z naszym zespołem wsparcia technicznego, który chętnie Ci pomoże.

Zespół AsWiseAI

