

• POD MASKĄ

Czego szuka **wyszukiwarka SumizAI**?

Wyszukiwarka jest miejscem, w którym osobista baza wiedzy zdaje egzamin, bo to jedyna funkcja używana w złości. **SumizAI indeksuje całą notatkę** — łącznie z oryginalną rozmową — i zwraca fragment, który pasował, a nie nazwę pliku.

Aktualizacja: 2026-08-14 8 pytań 3 min czytania

Q. Które części notatki są przeszukiwane?

Cztery, z różnymi wagami: **tytuł**, **podsumowanie**, **treść** oraz **oryginalne pytanie i odpowiedź** zapisane jako materiał źródłowy.

Ta czwarta jest powodem, dla którego wyszukiwanie działa tu inaczej. Połowę razy nie pamiętasz, jak notatka się nazywa ani jak została podsumowana — pamiętasz frazę z rozmowy, przykład, którego użył model, liczbę, którą podał. Ten tekst też jest w indeksie.

Wagi sprawiają, że trafienie w tytule nadal bije wzmiankę w głębi transkryptu, więc precyzja nie jest poświęcana dla zasięgu.

Q. Czy radzi sobie z odmianą i brakiem ogonków?

Z ogonkami tak i celowo. Tekst jest normalizowany przed indeksowaniem, a zapytanie normalizowane identycznym wyrażeniem — pisanie po polsku bez ogonków to normalne pisanie, nie błąd.

Ten szczegół z identycznym wyrażeniem nie jest pedanterią. Gdyby zapytanie normalizowało się inaczej niż indeks, baza po cichu przestałaby używać indeksu, a wyszukiwarka zrobiłaby się wolna, a nie błędna — czyli trudniej byłoby to zauważyć.

Odmiana jest obsługiwana dopasowaniem po przedrostku dla kandydatów i kontekstu, a nie słownikiem tematów. Używany słownik nie zna polskiej fleksji, więc *notatka* nigdy nie trafiłaby w *notatkami* — dopasowanie po przedrostku tę lukę zamyka.

Q. Dlaczego pokazywany jest fragment, a nie sama notatka?

Bo lista tytułów zmusza Cię do otwarcia pięciu notatek, żeby ustalić, którą miałeś na myśli.

Wynik niesie pasujący fragment z podświetloną frazą, więc już z listy widzisz, czy to ta notatka, w której wyliczyłeś model kosztów, czy ta, w której wspomniałeś o nim mimochodem.

Mówi też coś o strukturze notatki — trafienie w materiale źródłowym czyta się inaczej niż trafienie w podsumowaniu, a widzisz, które dostałeś.

Q. Czy mogę szukać w samym spisie treści?

Tak, i działa to celowo inaczej niż wyszukiwanie na liście. Nad drzewem rozdziałów jest pole wyszukiwania, a wpisanie w nie frazy **zwęża drzewo do trafionych gałęzi**.

Dostajesz więc strukturę zamiast płaskiej listy: widzisz, że cztery trafienia siedzą pod jednym rozdziałem, a jedno gdzieś zupełnie indziej — co często jest właściwą odpowiedzią na to, czego szukałeś.

Z notatki i z listy notatek prowadzi też odnośnik prosto do miejsca notatki w spisie treści, więc z trafienia przechodzisz do jego okolicy w jednym kroku.

Q. Jak wyszukiwanie ma się do tego, co widzi AI?

Blisko. Kiedy zadajesz pytanie w vaulcie, to ta sama maszyna znajduje najtrafniejsze notatki.

Twoje pytanie idzie razem ze spisem treści i pełną treścią maksymalnie pięciu notatek, w budżecie 24 000 znaków. Te pięć wybierane jest trafnością wobec pytania.

Jest jedna istotna różnica w budowie zapytań. Precyzyjne wyszukiwanie łączy Twoje słowa operatorem *i*, bo miałeś na myśli wszystkie. Szukanie kandydatów na podstawie długiego tekstu używa *lub*, bo wymaganie wystąpienia każdego słowa nie zwróciłoby nic.

Q. Co, jeśli wyszukiwarka przestaje znajdować notatki, o których wiem, że istnieją?

Przeindeksuj. Indeks to dane pochodne — kolumny wyliczone z treści notatek — a autorytatywne są pliki.

Operacja przeindeksowania odbudowuje indeks z zawartości magazynu. Jeśli indeks i pliki się nie zgadzają, wygrywa plik; nie ma scenariusza, w którym przeindeksowanie gubi notatkę.

Jest jedna ścieżka kodu, która zapisuje indeksowaną treść, dokładnie po to, żeby żaden zapis nie zaktualizował trzech z czterech przeszukiwanych kolumn i nie zostawił po cichu czwartej nieaktualnej.

Q. Czy wyszukiwanie jest ograniczone do vaulta?

Tak, jak wszystko inne. Vault jest granicą dla wyszukiwania dokładnie tak samo jak dla kontekstu modelu.

Jest to wymuszone strukturalnie, a nie pamiętaniem o filtrze: repozytoria nie wystawiają sposobu na odpytanie notatek bez właściciela i vaulta.

Praktyczny efekt jest taki, że wyszukiwanie pozostaje celne, gdy rośnie łączna liczba notatek, bo zawsze przeszukujesz jedną dziedzinę, a nie wszystko, co kiedykolwiek napisałeś.

Q. Czego wyszukiwarka świadomie nie robi?

Nie jest semantyczna. Nie ma embeddingów ani bazy wektorowej; wyszukiwanie jest leksykalne, z podobieństwem trigramowym do łapania bliskich chybień.

Dla tej aplikacji to właściwy wybór. Krok semantyczny odbywa się tam, gdzie jest najcenniejszy — model czyta wybraną krótką listę i ją ocenia — zamiast rozmazywać się po indeksie, który trzeba utrzymywać, migrować i przeliczać za każdą zmianą modelu.

Utrzymuje to też całość szybką i sprawdzalną. Widzisz, dlaczego wynik pasował.