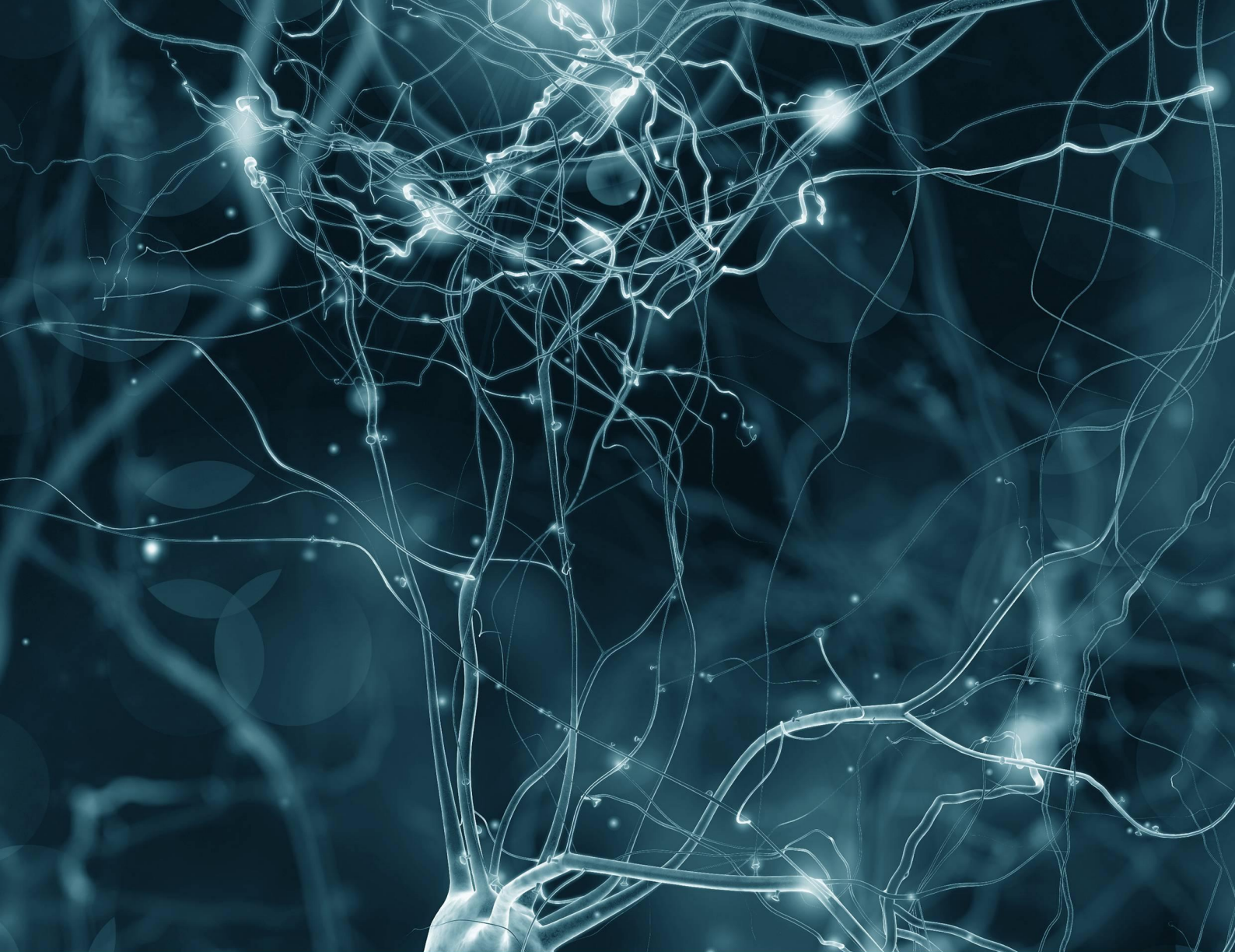




AGiDOT

Sztuczna inteligencja w biznesie

Stawiamy na transformację biznesu poprzez wykorzystanie sztucznej inteligencji

RAG na danych firmowych

Zaufanie, prywatność i modele lokalne — co AI naprawdę potrafi na Twoich dokumentach

Coraz więcej firm chce wykorzystać sztuczną inteligencję do pracy na własnych dokumentach — umowach, procedurach, ofertach, danych kadrowych. Powstrzymują je zwykle dwie obawy: „**AI wszystko zmyśla, nie można jej ufać**” oraz „**to zagrożenie — wyciekną moje dane**”.

Ten raport odpowiada na obie — nie deklaracjami, lecz liczbami. Na danych fikcyjnej firmy budowlanej przeprowadziliśmy **294 kontrolowane testy**, porównując modele lokalne z chmurowymi, różne strategie wyszukiwania oraz metody obrony przed halucynacją. Pokazujemy, co da się dziś zrobić na własnej infrastrukturze, a czego jeszcze nie.

Spis treści

Wprowadzenie	5
Abstrakt	5
Dlaczego to ważne	5
Dla kogo jest ten dokument	5
Co znajdziesz w środku	6
RAG kontra zwykły chatbot	7
Skąd model bierze odpowiedź	7
Różnica widoczna w liczbach	7
Pełna kontrola nad wyszukiwaniem	7
Obawa pierwsza: „AI wszystko zmyśla”	9
Halucynacja to problem inżynierski, nie wyrok	9
Dowód z testów	9
Test, którego celem było przyłapanie na kłamstwie	10
Obawa druga: „Wyciekną moje dane”	11
Dane nie muszą opuszczać firmy	11
Lokalny model wystarcza	11
Warstwa bezpieczeństwa niezależna od modelu	11
Uczciwie o kompromisach	12
Który model lokalny wybrać	13
Rozmiar nie decyduje o jakości	13
Jakość ma swoją cenę: czas i koszt	14
Co to znaczy w praktyce	14
Co robić, gdy plik jest duży	15
Problem: rzadki fakt w tysiącach stron	15
Samo wyszukiwanie znaczeniowe nie wystarcza	15
Długi kontekst pomaga, ale się nie skaluje	15
Znalezienie to nie wszystko — trzeba jeszcze wydobyć	16
Zalecenia dla dużych plików	16
Jak testowaliśmy	17
Fikcyjna firma, realne pułapki	17
Trzy testy, 294 pomiary	17
Jak ocenialiśmy poprawność	17

Wnioski i rekomendacje	19
Co pokazały testy	19
Rekomendowany punkt startowy	19
Granice, o których mówimy wprost	19
Kiedy RAG nie jest potrzebny	20
AGIDOT — AI w praktyce	21

Wprowadzenie

Abstrakt

Dwie obawy najczęściej powstrzymują firmy przed wdrożeniem AI na własnych danych: przekonanie, że model „wszystko zmyśla”, oraz strach o wyciek poufnych informacji. Oba są uzasadnione — ale oba mają dziś konkretną, techniczną odpowiedź. Tą odpowiedzią jest **RAG** (Retrieval-Augmented Generation): architektura, w której model odpowiada nie z „wiedzy z internetu”, lecz wyłącznie na podstawie wskazanych, firmowych dokumentów — z możliwością sprawdzenia źródła każdej odpowiedzi.

Aby nie poprzestać na deklaracjach, zbudowaliśmy realny system RAG i przetestowaliśmy go na danych fikcyjnej firmy budowlanej „Smart-Bud”. Zadaliśmy mu 12 pytań o różnym stopniu trudności i sprawdziliśmy, jak radzą sobie z nimi różne modele (lokalne i chmurowe) oraz różne strategie wyszukiwania. Łącznie wykonaliśmy **294 pomiary**. Ten raport przedstawia ich wyniki i płynące z nich wnioski.

Dlaczego to ważne

Wiedza firmy jest rozproszona — w dziesiątkach dokumentów, tabel i procedur, których nikt nie pamięta w całości. Zwykły chatbot (jak publiczny asystent w przeglądarce) tej wiedzy **nie zna** i, zapytany o firmowy szczegół, albo się przyzna do niewiedzy, albo zmyśli prawdopodobnie brzmiącą odpowiedź. To właśnie źródło obawy „AI zmyśla”.

Jednocześnie wysłanie poufnych dokumentów do zewnętrznego modelu w chmurze bywa nie do zaakceptowania — ze względu na tajemnicę handlową, dane osobowe czy wymogi RODO. To źródło obawy „dane wyciekną”.

RAG odpowiada na oba problemy naraz: **osadza** odpowiedzi w firmowych źródłach (ograniczając zmyślanie) i może działać **w całości lokalnie** (bez wysyłania danych na zewnątrz). Pytanie brzmi już nie „czy”, lecz „jak dobrze” — i to jest przedmiotem tego raportu.

Dla kogo jest ten dokument

To dokument dla **decydentów w firmach**, które rozważają wykorzystanie AI na własnych dokumentach i chcą zrozumieć realne możliwości oraz granice tej technologii — bez marketingowego uproszczenia. Część techniczna (rozdziały o modelach i strategiach wyszukiwania) przyda się również zespołom IT oceniającym wdrożenie u siebie.

Co znajdziesz w środku

- **Czym RAG różni się od zwykłego chatbota** i dlaczego to różnica fundamentalna, a nie kosmetyczna.
- **Obrona przed halucynacją** — jak, warstwami, ograniczyć zmyślanie, poparte wynikami testów.
- **Modele lokalne** — czy własna infrastruktura wystarcza i który model wybrać (odpowiedź zaskakuje: rozmiar nie decyduje).
- **Duże pliki** — co robić, gdy dokument ma tysiące stron i szukany fakt pojawia się tylko raz.
- **Metodyka i granice** — jak testowaliśmy oraz gdzie RAG jeszcze dziś zawodzi.

RAG kontra zwykły chatbot

Skąd model bierze odpowiedź

Zwykły chatbot odpowiada z wiedzy „zapamiętanej” podczas treningu. O firmie Smart-Bud, jej pracownikach czy cenniku nie wie nic — bo tych danych nigdy nie widział. Zapytany, wygeneruje odpowiedź statystycznie prawdopodobną, ale niezwiązaną z rzeczywistością Twojej firmy.

System RAG działa inaczej. Zanim model odpowie, **wyszukiwarka** odnajduje w firmowych dokumentach fragmenty najbardziej związane z pytaniem i **wkłada je do polecenia** dla modelu. Model dostaje więc nie tylko pytanie, lecz i konkretne wyciągi ze źródeł — i ma polecenie odpowiadać wyłącznie na ich podstawie.

	Zwykły chatbot	RAG na danych firmy
Źródło odpowiedzi	wiedza modelu (nieznana)	dokumenty firmy (znane, cytowane)
Aktualność	do daty treningu	zawsze aktualne dokumenty
Sprawdzalność	brak	cytat + wskazanie pliku
Dane wrażliwe	wysyłane do modelu	mogą zostać lokalnie
Zmyślanie	trudne do kontroli	ograniczone osadzeniem w źródłach

Różnica widoczna w liczbach

Ta różnica nie jest teoretyczna. W naszych testach ten sam model, zapytany o firmowe fakty **bez** dostępu do dokumentów (tryb chatbota), odpowiadał poprawnie tylko w **8–17%** przypadków. Po podłączeniu tych samych pytań do RAG poprawność rosła do **58–92%** — zależnie od modelu i strategii wyszukiwania. Szczegóły w kolejnym rozdziale.

Pełna kontrola nad wyszukiwaniem

Najważniejsza przewaga RAG nad zamkniętym chatbotem jest często pomijana: to **my kontrolujemy, co model dostaje do przeczytania**. Możemy wymusić uwzględnienie dokumentów krytycznych, filtrować wyniki po uprawnieniach użytkownika (kto ma prawo zobaczyć dany dokument), łączyć wyszukiwanie znaczeniowe z dokładnym

dopasowaniem słów i gwarantować, że pewne kategorie dokumentów zawsze zostaną przeszukane.

Ta kontrola ma jednak granice, o których trzeba wiedzieć — nie da się po prostu „wczytać wszystkiego” przy każdym pytaniu, bo duże zbiory przekraczają możliwości modelu i topią właściwą odpowiedź w szumie. Jak sobie z tym radzić, pokazujemy w rozdziale o dużych plikach.

Obawa pierwsza: „AI wszystko zmyśla”

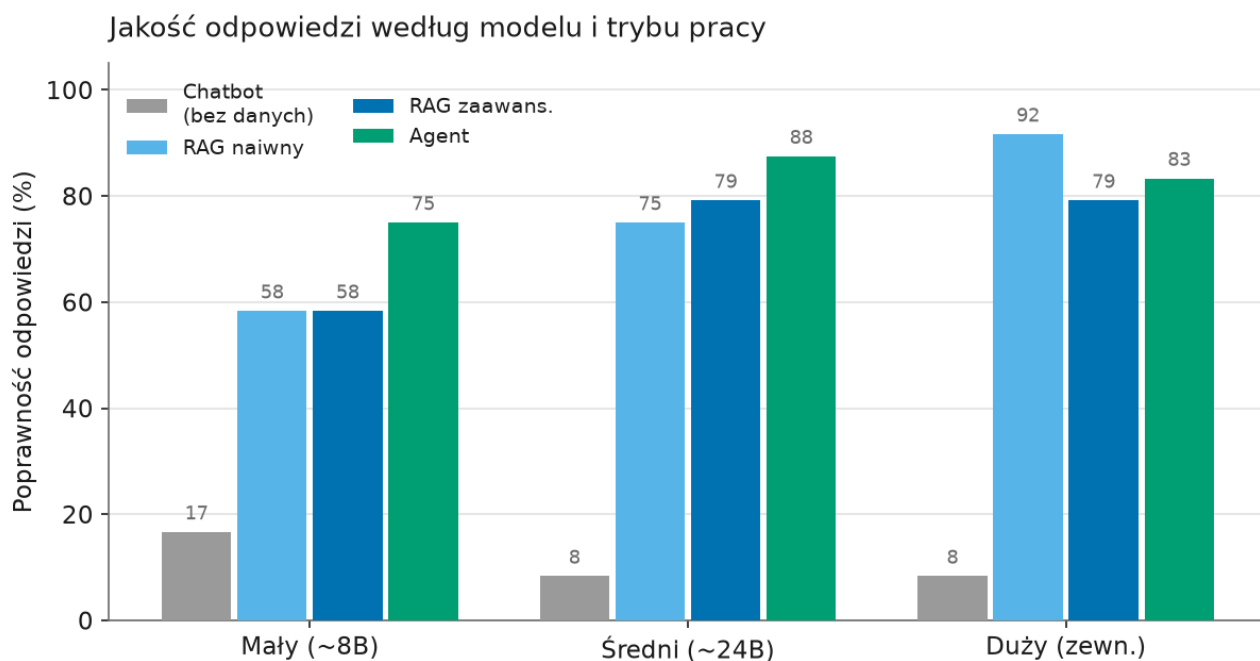
Halucynacja to problem inżynierski, nie wyrok

Halucynacja — pewnie brzmiąca, ale nieprawdziwa odpowiedź — jest realnym zjawiskiem. Nie jest jednak cechą, z którą trzeba się pogodzić. To problem, który **redukuje się warstwami**, a każda warstwa domyka kolejne źródło błędu.

1. **Osadzenie w źródłach.** Model odpowiada wyłącznie z dostarczonych fragmentów i wskazuje, z którego pliku pochodzi fakt. Odpowiedź bez pokrycia w źródłach to sygnał ostrzegawczy.
2. **Prawo do „nie wiem”.** Gdy w dokumentach nie ma odpowiedzi, model ma ją tak zakomunikować — zamiast zmyślać. To kluczowa, często pomijana warstwa.
3. **Jakość wyszukiwania.** Lepsze wyszukiwanie i porządkowanie fragmentów to mniej okazji do halucynacji — śmieci na wejściu dają śmieci na wyjściu.
4. **Weryfikacja odpowiedzi.** Drugi przebieg (lub druga instancja modelu) sprawdza, czy każde zdanie odpowiedzi wynika ze źródeł.
5. **Człowiek w pętli.** Przy decyzjach krytycznych odpowiedź AI jest propozycją do zatwierdzenia, nie wyrokiem.

Dowód z testów

Sprawdziliśmy dwie skrajności na tych samych pytaniach. W trybie chatbota — bez dostępu do firmowych dokumentów — model musi zgadywać; poprawność spada do kilkunastu procent, a „poprawne” trafienia dotyczą głównie wiedzy ogólnej lub uczciwego przyznania się do niewiedzy. Po włączeniu RAG z poleceniem „odpowiadaj tylko ze źródeł” jakość rośnie wielokrotnie.



Wykres pokazuje trzy klasy modeli w czterech trybach pracy. Wnioski:

- **Sam chatbot (szare słupki) zawodzi** przy pytaniach o dane firmy — 8–17% poprawności. To dokładnie zjawisko, które użytkownik nazywa „AI zmyśla”.
- **RAG podnosi jakość dramatycznie** — do 58–92%. To ta sama AI, ale z dostępem do zweryfikowanych źródeł.
- **Tryb agentowy** (model sam dopytuje o kolejne dokumenty) dodatkowo pomaga słabszym modelom — mały model 8B rośnie z 58% do 75%.

Test, którego celem było przyłapanie na kłamstwie

Jedno z pytań celowo dotyczyło informacji, której **w dokumentach nie ma** (marża firmy na konkretnym projekcie). Dobrze zbudowany system RAG odpowiada tu „brak danych w dokumentach”, a nie zmyśla liczbę. W tym scenariuszu poprawnie zachowały się wszystkie testowane modele — potwierdzając, że powściągliwość da się wymusić projektem systemu, a nie tylko nadzieją. To najważniejsza pojedyncza odpowiedź na obawę o zaufanie: **system można zbudować tak, by wolał powiedzieć „nie wiem” niż zmyślić.**

Obawa druga: „Wyciekną moje dane”

Dane nie muszą opuszczać firmy

Najprostsza odpowiedź na obawę o wyciek: **nie wysyłaj danych na zewnątrz**. System RAG można zbudować w całości na infrastrukturze firmy (lub na dedykowanym, odseparowanym serwerze), tak że dokumenty i pytania nigdy nie trafiają do zewnętrznego dostawcy.

Składają się na to dwa elementy, które da się uruchomić lokalnie:

- **Model wyszukiwający (embeddingi)**. Zamienia dokumenty i pytania na reprezentacje liczbowe, po których działa wyszukiwarka. Dobre modele wielojęzyczne (z obsługą polskiego) są niewielkie i uruchamialne na zwykłym serwerze.
- **Model językowy (LLM)**. Generuje odpowiedź z dostarczonych fragmentów. Tu leży sedno pytania: czy model, który *da się* uruchomić lokalnie, jest wystarczająco dobry?

Lokalny model wystarcza

Odpowiedź brzmi: tak — i to jest najważniejszy wniosek tego rozdziału. W naszych testach **średni model open-weight (~24B parametrów)**, możliwy do uruchomienia na jednej mocniejszej karcie graficznej, w trybie agentowym osiągnął **88%** poprawności — praktycznie na poziomie dużego modelu chmurowego (83–92%). Różnica jakości, za którą trzeba płacić utratą kontroli nad danymi, okazuje się niewielka.

To zmienia charakter decyzji. Nie chodzi już o wybór „prywatność albo jakość” — dla większości zastosowań firmowych można mieć oba naraz.

Warstwa bezpieczeństwa niezależna od modelu

Prywatność to nie tylko „gdzie stoi model”. Architektura RAG pozwala dołożyć zabezpieczenia, których zamknięty chatbot nie oferuje:

- **Kontrola dostępu per dokument** — użytkownik dostaje odpowiedzi tylko z dokumentów, do których ma uprawnienia.
- **Ślad audytowy** — kto, o co pytał i na jakich źródłach oparto odpowiedź.
- **Odseparowanie od sieci publicznej** — wdrożenie może działać bez dostępu do internetu.

-
- **Zgodność z RODO** — dane osobowe i tajemnica handlowa zostają w granicach firmy.

Uczciwie o kompromisach

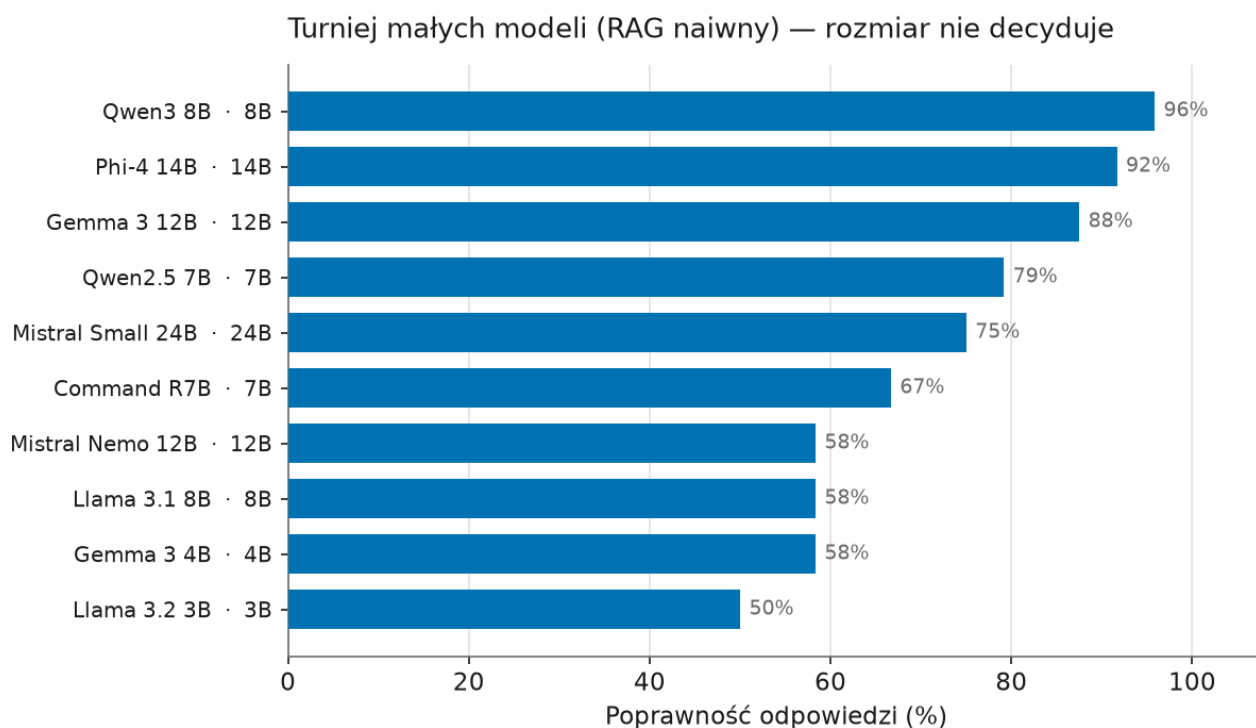
Model lokalny bywa słabszy od absolutnie najlepszego modelu chmurowego w najtrudniejszych zadaniach rozumowania. Dlatego sensownym wariantem bywa **podejście hybrydowe**: dane wrażliwe obsługuje model lokalny, a zadania nieobejmujące poufnych treści (np. przeredagowanie tekstu) mogą trafić do mocniejszego modelu w chmurze. Granicę ustala się świadomie — to decyzja projektowa, nie przypadek.

Który model lokalny wybrać

Rozmiar nie decyduje o jakości

Skoro model może działać lokalnie, powstaje pytanie praktyczne: który? Intuicja podpowiada „największy, na jaki nas stać”. Testy pokazują, że to intuicja myląca.

Przepuściliśmy te same 12 pytań (w trybie RAG) przez **10 różnych modeli open-weight** do ~24B parametrów — czyli takich, które realnie uruchomisz na własnym sprzęcie. Wyniki układają się wbrew intuicji o rozmiarze.



- **Największy model w stawce (24B) zajął dopiero 5. miejsce** (75%), przegrywając z modelami 8B, 12B i 14B.
- **Dwa modele tej samej wielkości (8B) dzieli przepaść** — 96% kontra 58%. O jakości decyduje rodzina i sposób treningu, nie liczba parametrów.
- **Nawet model 3B osiąga 50%** — wystarcza do prostych wyszukikań, choć zawodzi w rozumowaniu.

Wniosek dla firmy jest praktyczny i oszczędny: **nie trzeba kupować najcięższego sprzętu**. Trzeba dobrać właściwy model.

Jakość ma swoją cenę: czas i koszt

Zwycięzca rankingu (model 8B, 96%) to model „rozumujący” — myśli „na głos”, zanim odpowie. Płaci za to czasem (~10 sekund na odpowiedź) i liczbą przetwarzanych tokenów. Dla wielu zastosowań to zła wymiana.

Rola	Rekomendacja	Dlaczego
Najwyższa jakość	model „rozumujący” 8B	96%, ale wolny i kosztowny w obliczeniach
Najlepszy kompromis	model 12–14B (np. klasy Gemma / Phi)	88–92% przy 2–3 s i zwięzłych odpowiedziach
Najlżejszy sprzęt	model 3–4B	50–58%, do prostych wyszukiwań i lookupów

Dla typowego wdrożenia firmowego rekomendujemy **model klasy 12–14B** jako punkt startowy: bardzo dobra jakość, szybkie i zwięzłe odpowiedzi, umiarkowane wymagania sprzętowe.

Co to znaczy w praktyce

Dobór modelu to nie zakup „najdroższego pudełka”, lecz dopasowanie do zadania i sprzętu. Ponieważ modele open-weight zmieniają się co kilka tygodni, właściwym podejściem jest **mierzenie na własnych danych** — dokładnie tak, jak zrobiliśmy to w tym raporcie — zamiast opierania się na rankingach ogólnych.

Co robić, gdy plik jest duży

Problem: rzadki fakt w tysiącach stron

Firmowe dokumenty potrafią być ogromne — wielostronicowe regulaminy, raporty, archiwa. Kluczowy fakt (numer, kod, jedno zdanie procedury) często pojawia się w nich **tylko raz**. Czy system go znajdzie?

Sprawdziliśmy to na parze dokumentów po **~2,5 MB każdy** (setki tysięcy słów), w których zaszyliśmy trzy fakty o malejącej częstotliwości: fakt A — 3 razy na plik, fakt B — 2 razy, fakt C — tylko raz. Następnie porównaliśmy strategie wyszukiwania. Wyniki są pouczające.

Samo wyszukiwanie znaczeniowe nie wystarcza

Najczęściej stosowane wyszukiwanie **znaczeniowe** (po sensie, przez embeddingi) **zgubiło rzadkie fakty**. Fakt występujący raz na tysiące fragmentów nie trafił do wyników — a zwiększenie liczby pobieranych fragmentów nic nie dało, bo problemem nie była liczba, lecz brak dopasowania znaczeniowego. Dla dosłownych faktów (kody, liczby, nazwy) ratunkiem okazało się **wyszukiwanie po słowach kluczowych (BM25)**, które odnalazło wszystkie trzy fakty.

Strategia	Fakt A (3×)	Fakt B (2×)	Fakt C (1×)
Wyszukiwanie znaczeniowe	znajduje	gubi	gubi
Słowa kluczowe (BM25)	znajduje	znajduje	znajduje
Hybryda (znaczenie + słowa)	znajduje	znajduje	znajduje
Cały plik do modelu 1M	znajduje	znajduje	znajduje*

* Dla pojedynczego pliku. Przy połączeniu obu plików dokument przekroczył okno kontekstu i wariant „cały plik” przestał być wykonalny.

Długi kontekst pomaga, ale się nie skaluje

Kuszące jest „wrzucić cały plik do modelu o dużym oknie kontekstu” (np. 1 milion tokenów). Dla **jednego** pliku 2,5 MB zadziałało — model znalazł nawet fakt występujący raz. Ale po połączeniu dwóch takich plików dokument **przekroczył okno kontekstu** i strategia stała się fizycznie niewykonalna. Długi kontekst to użyteczne narzędzie do pojedynczego dużego dokumentu — ale **nie skaluje się** na rosnące archiwum. RAG skaluje się.

Znalezienie to nie wszystko — trzeba jeszcze wydobyć

Warto rozróżnić dwa etapy. Nawet gdy wyszukiwarka *dostarczyła* fakt do kontekstu, słabszy model potrafił go **nie wydobyć** z odpowiedzi. Sukces wymaga obu: dobrego wyszukiwania **i** wystarczająco dobrego modelu do wyciągnięcia odpowiedzi.

Zalecenia dla dużych plików

- **Nie polegaj na samych embeddingach.** Minimum to hybryda: wyszukiwanie znaczeniowe **plus** słowa kluczowe, z ponownym porządkowaniem wyników.
- **Dla dosłownych faktów** (numery umów, kody, NIP-y, nazwy własne) dołącz dokładne wyszukiwanie — embeddingi je gubią.
- **Większa liczba pobieranych fragmentów to nie lekarstwo** na brak dopasowania — dokłada szum, nie trafność.
- **Długi kontekst** stosuj świadomie: dobry do pojedynczego dużego pliku, bezradny wobec dużego archiwum.
- **Krytyczne fakty warto wzmocnić w indeksie** (metadane, streszczenia, osobne pola) — pojedyncze wystąpienie łatwo ginie.

Jak testowaliśmy

Fikcyjna firma, realne pułapki

Aby test był rzetelny i powtarzalny, stworzyliśmy dane fikcyjnej firmy budowlanej „**Smart-Bud**”: 17 dokumentów w ośmiu działach — dane firmy, pracownicy i ich uprawnienia, projekty, cenniki, sprzęt, dostawcy, procedury BHP oraz finanse. Dokumenty zaprojektowaliśmy tak, by zawierały realistyczne **pułapki**, na których widać różnicę między podejściami:

- **Fakty rozproszone** — odpowiedź wymaga połączenia kilku dokumentów (np. „czy mamy dostępnego operatora żurawia w danym tygodniu” wymaga zestawienia uprawnień z kalendarzem urlopów).
- **Wersjonowanie** — dwa cenniki (stary i aktualny); poprawna odpowiedź to cena z aktualnego.
- **Arytmetyka z warunkiem** — suma wartości aktywnych kontraktów, z pominięciem kontraktu zakończonego.
- **Igła w stogu siana** — pojedynczy fakt ukryty w długiej procedurze.
- **Pytanie bez odpowiedzi** — informacja, której w danych nie ma (test obrony przed halucynacją).

Trzy testy, 294 pomiary

Na tych danych przeprowadziliśmy trzy niezależne testy:

Test	Zakres	Pomiary
1. Modele i tryby	3 klasy modeli × 4 tryby pracy × 12 pytań	144
2. Turniej modeli	10 modeli open-weight ≤24B, tryb RAG	120
3. Duże pliki	5 strategii × 3 fakty × 2 warianty	30

Jak ocenialiśmy poprawność

Odpowiedzi oceniano dwustopniowo. Tam, gdzie liczył się dokładny fakt (kod, liczba), stosowaliśmy **sprawdzenie automatyczne** — obiektywne i tanie. Dla odpowiedzi opisowych w języku polskim rozstrzygał **model- sędzia**: mocny, niezależny model porównujący odpowiedź ze wzorcem (prawdą) według jasnych kryteriów, z rozróżnieniem odpowiedzi pełnej, częściowej i błędnej. Odmowa odpowiedzi na pytanie faktograficzne liczyła się jako błąd — inaczej „nie wiem” fałszywie zawyżałoby wyniki.

Cały eksperyment jest **powtarzalny**: generacja jest deterministyczna, a wyniki buforowane. Ten sam zestaw testów można uruchomić ponownie na danych konkretnej firmy — i to właśnie rekomendujemy przed każdym poważnym wdrożeniem.

Wnioski i rekomendacje

Co pokazały testy

- **RAG odpowiada na obawę „AI zmyśla”.** Osadzenie odpowiedzi w firmowych źródłach podnosi poprawność z kilkunastu procent (sam chatbot) do 58–92%. System można zbudować tak, by wolał przyznać „nie wiem” niż zmyślić.
- **RAG odpowiada na obawę „dane wyciekną”.** Lokalny model ~24B dorównuje modelowi chmurowemu, a dane nie muszą opuszczać firmy.
- **Dobór modelu jest ważniejszy niż jego rozmiar.** Model 12–14B bywa lepszy niż 24B — i szybszy. Nie trzeba najcięższego sprzętu.
- **Przy dużych plikach liczy się strategia wyszukiwania.** Same embeddingi gubią rzadkie fakty; hybryda ze słowami kluczowymi jest obowiązkowa, a długi kontekst nie skaluje się na archiwum.

Rekomendowany punkt startowy

Dla typowego wdrożenia firmowego proponujemy:

- **model lokalny klasy 12–14B** (open-weight, uruchamialny u siebie),
- **wyszukiwanie hybrydowe** (znaczeniowe + słowa kluczowe) z ponownym porządkowaniem wyników,
- **osadzanie odpowiedzi w źródłach** z cytatami i prawem do „nie wiem”,
- **kontrolę dostępu i ślad audytowy** od pierwszego dnia,
- **pilotaż na jednym zbiorze danych z pomiarem jakości** przed rozszerzeniem — tak jak w tym raporcie.

Granice, o których mówimy wprost

RAG nie jest magią. W naszych testach **wszystkie** modele zawiodły w dwóch typach zadań:

- **złożona arytmetyka z warunkami** (np. suma wartości z pominięciem pozycji o określonym statusie),
- **wieloetapowe rozumowanie o dostępności zasobów** (zestawienie sprzętu, ludzi i terminów w jeden wniosek).

To zadania, w których „gołe” wyszukiwanie nie wystarcza — potrzebne są dodatkowe narzędzia (np. liczenie po stronie systemu, a nie modelu) i weryfikacja. Znajomość tych granic jest częścią odpowiedzialnego wdrożenia: wiemy, co dziś działa, i wiemy, gdzie dołożyć zabezpieczenia.

Kiedy RAG nie jest potrzebny

Dla uczciwości: jeśli zadanie nie dotyczy firmowych danych — np. przeredagowanie tekstu czy burza mózgów — RAG nie jest potrzebny. Jego wartość ujawnia się dokładnie tam, gdzie liczy się **Twoja wiedza**: aktualna, poufna i sprawdzalna.

AGIDOT — AI w praktyce

AGIDOT to firma ekspercka od praktycznego zastosowania sztucznej inteligencji w biznesie. Nie sprzedajemy hasła „AI” — pokazujemy, jak realnie odciąża ludzi z powtarzalnej pracy. Rozwiązania, które budujemy w projektach takich jak ten, oferujemy też jako gotowe produkty i wdrożenia.

SZKOLENIA

Szkolenia z AI

Uczymy zespoły pracować z agentami AI na co dzień — od promptów po automatyzację własnych narzędzi. Bez teorii oderwanej od praktyki.

WDROŻENIA

Wdrożenia AI

Projektujemy i wdrażamy asystentów AI, systemy RAG na dokumentach firmowych i automatyzację obiegu dokumentów — zintegrowane z Waszymi systemami.

NARZĘDZIA

Gotowe narzędzia

Sprawdzone w naszych projektach narzędzia dostępne od ręki w sklepie — z kluczem licencyjnym i wsparciem.

Przykłady ze sklepu

Kilka produktów i usług z hub.agidot.eu/sklep:



SZKOLENIE ONLINE

Praca z AI w codziennych zadaniach — szkolenie online

Nagrania wideo pokazujące praktyczną pracę z agentami AI: od prostych promptów po automatyzację własnych narzędzi.

999 zł

[Szczegóły →](#)



USŁUGA

Wdrożenie systemu RAG w firmie

Kompleksowe wdrożenie systemu RAG na dokumentach firmowych: analiza, instalacja, integracja i szkolenie zespołu. Wycena indywidualna.

wycena indywidualna

[Szczegóły →](#)

Zobacz cały sklep → hub.agidot.eu/sklep

Z naszego bloga

Ostatnie wpisy — konkretnie o AI w firmie, bez lania wody. Więcej na agidot.eu/blog:

13 LIPCA 2026

5 procesów dokumentowych, które AI może przejąć już dziś

Od faktur zakupowych po korespondencję — pięć miejsc, w których sztuczna inteligencja najszybciej przejmuje pracę z dokumentami i najszybciej się zwraca.

6 LIPCA 2026

Od skanu do systemu: jak AI automatyzuje obieg dokumentów w firmie

Krok po kroku droga faktury: od PDF-a w skrzynce, przez OCR i walidację, po gotowy wpis w systemie księgowym — i ile to daje w liczbach.

Newsletter

Raz w tygodniu, w poniedziałek rano, jedna konkretna dawka wiedzy o tym, jak AI realnie pomaga w firmie — bez marketingowego szumu i bez obiecywania cudów.

- jeden temat tygodnia rozwinięty na blogu,
- praktyczna wskazówka do wdrożenia od ręki,
- zapowiedź materiałów z całego tygodnia.

Zapisz się → hub.agidot.eu/newsletter