

AI in academic research - practical workshop

Bielik na Twoim laptopie — prywatna asystentka badawcza

Maksymilian Małecki · Tech Lead @ Insly · certyfikowany trener Eskadry Bielika
WPiA UAM · 28 czerwca 2026

Kto mówi

Maksymilian Małecki

- **Tech Lead @ Insly** — produkcyjny RAG (**OWU AI**) na poliglotycznej platformie insurtech
- **17+ lat w IT** · wcześniej NordSecurity, Bitnoi.se, ORO, Arena.pl
- **Na scenie:** 3× PHPers Summit, DevAI 2025, Eskadra Bielika
- **Dziś:** AI w pracy naukowej, ale na własnym laptopie i bez wysyłania niczego na zewnątrz

Buduję RAG-i z polskim językiem prawniczym na produkcji. Wiem, co działa.

Krok 0 — pobierz LM Studio teraz

Zanim wejdziemy w treść — pobierz LM Studio na swój laptop.

Darmowe, działa offline, Windows / Mac / Linux. Potrzebujesz **6 GB wolnego miejsca na dysku.**

1. Wejdź na **lmstudio.ai** lub zeskanuj QR
2. Pobierz wersję dla swojego systemu
3. Zainstaluj — standardowy instalator

Zajmie 2–3 minuty. Zrób to **teraz** — model pobierzemy razem za chwilę.



lmstudio.ai

Plan na 90 minut

Czas	Akt	Co	Twoja rola
0:05–0:18	I	Krajobraz płatnych AI (ChatGPT, NotebookLM, ElevenLabs STT, Gamma)	patrzysz, ładujesz model
0:18–0:25	II	Ryzyka tych narzędzi (GDPR, koszt, lock-in, halucynacje)	słuchasz
0:25–0:32	II → III	Mapa: płatne → open-source + kompromisy	sanity check
0:32–0:36	III	Jak rozmawiać z Bieliem + załaduj model	słuchasz, ładujesz
0:36–1:03	III	Hands-on: czyść dane — OCR, transkrypcja, anonimizacja	pracujesz
1:03–1:16	III	Demo: provider-switch + transkrypcja audio	patrzysz
1:16–1:30	—	Open mic + materiały + ankieta	dyskutujesz

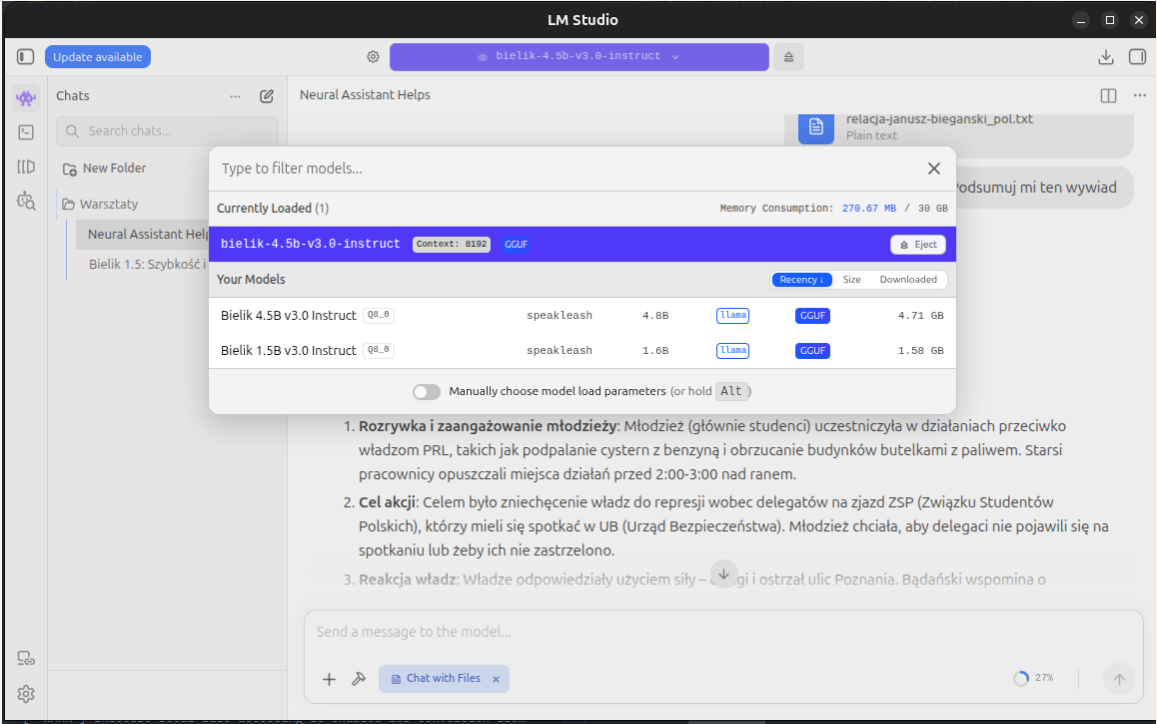
Który model Bielika pobrać?

W LM Studio → **Discover** → wyszukaj **bielik**

RAM laptopa	Model do pobrania	Rozmiar
16 GB lub więcej	Bielik 4.5B v3.0 Instruct Q8_0	~4,7 GB
Mniej niż 16 GB	Bielik 1.5B v3.0 Instruct Q8_0	~1,6 GB

Nie wiesz ile masz RAM? → Windows: **Ctrl+Shift+Esc**
→ Wydajność → Pamięć · Mac: Apple → O tym Macu

Pobierz **teraz** — zajmie kilka minut. Model będzie gotowy gdy wejdziemy w hands-on.



Wszystko zostaje w laboratorium.

Nic nie wychodzi do BigTechów.

AKT I

Co potrafią płatne AI?

Najważniejsze narzędzia AI dla doktoranta — *przegląd*

Kategoria	Lider płatny	Koszt/mies
Chat ogólny	ChatGPT Plus / Claude Pro	20 USD
RAG nad PDF-ami ★	NotebookLM (Google)	0–20 USD
Wyszukiwarka naukowa	Elicit / scite.ai / Consensus	10–75 USD
Web research z LLM	Perplexity Pro	20 USD
Tłumaczenie	DeepL Pro	9 USD
Pisanie / styl	Grammarly / DeepL Write	12–30 USD
Transkrypcja audio	ElevenLabs Speech-to-Text	22 USD
OCR + struktura	Adobe Acrobat AI / ABBYY	25–30 USD
Second brain	Notion AI / Mem	20 USD
Generator slajdów	Gamma / Beautiful.ai	12–25 USD

Suma pełnego stacku doktoranta: 130–170 USD/mies. ≈ 500–700 PLN/mies.

Pobierz starter-pack teraz

Zanim przejdziemy do demo — pobierz paczkę, żebyś mógł/mogła przejść razem ze mną przez płatne narzędzia.

<https://workshops.mmx3.pl/konferencja-krd-2026/starter-pack.zip>

Zawiera: instrukcję instalacji, gotowe prompty i korpus ćwiczeniowy.



Demo #1 — NotebookLM (Google)

Bezpośredni odpowiednik tego, co za chwilę zbudujemy lokalnie.

- Wrzucasz do 50 PDF-ów, Google Doców albo linków
- Pytasz po polsku, dostajesz odpowiedź z cytowaniami źródeł
- Killer-feature: generuje podcast z Twoich materiałów (dwie osoby dyskutują)
- Free tier wystarcza

Pytanie testowe na żywo: *"Streść mi wnioski z tych 3 artykułów o X. Co wspólne, co odmienne?"*

Co warto zapamiętać: Cała Twoja biblioteka — w systemie Google. Pamiętaj o tym.

Demo #2 — Perplexity / ChatGPT Search

Wyszukiwanie z LLM, z linkami do źródeł.

- „Researchuj mi temat X” — dostajesz strukturę plus linki do publikacji
- Tryb „Academic” filtruje do recenzowanych źródeł
- Bije Google Scholar w sytuacjach, gdy nie wiesz jeszcze, jakich słów kluczowych użyć

Pytanie testowe na żywo: *"Jakie są najnowsze prace o LLM-ach w analizie polskich aktów prawnych z lat 2024-2026?"*

Co warto zapamiętać: Świetne na **publiczne** info. Twoja rozprawa — nie idzie tam.

Demo #3 — ElevenLabs Speech-to-Text (transkrypcja)

Dla każdego, kto robi wywiady, etnografię albo focus groupy.

- Audio i wideo zamienia na transkrypcję w czasie rzeczywistym
- Diaryzacja, czyli kto co powiedział
- Generuje summary i tagi tematyczne
- Można pisać do nagrania jak do dokumentu

Pytanie: Ile masz godzin nagrań wywiadów? *(reka w górę)*

Co warto zapamiętać: Twoi respondenci podpisali zgodę na **ElevenLabs w USA**? Bo to tam idzie ich głos.

Demo #4 — Gamma (generator slajdów)

Tekst zamienia na estetyczną prezentację w 30 sekund.

- Wklejasz tytuł i bullet pointy, dostajesz kompletny deck
- Sprawdza się na seminaria, obrony, konferencje wewnętrzne
- Eksport do PPTX, PDF albo Google Slides

Co warto zapamiętać: Twój *unpublished* draft pracy, wrzucony do Gammy, żeby zrobić slajdy, **jest u nich**.



gamma.app

Krajobraz w jednym zdaniu

Płatne AI robi rzeczy, których pięć lat temu nie było — tłumaczenie, RAG, transkrypcja, generowanie slajdów. Jednym kliknięciem, w jakości godnej publikacji.

Ale każde z tych narzędzi to jednocześnie trzy rzeczy: dane wysyłane do BigTechu w USA albo Chinach, subskrypcja sumująca się do ~700 zł miesięcznie i lock-in w cudze API.

Czas porozmawiać o kosztach tego stacku.

Sprawdź setup — czy masz Bielika?

Otwórz LM Studio →  **Chat** → „Select a model to load”.

- **Widzisz** `Bielik 4.5B v3.0 Instruct` ? (pobrany w domu) → załaduj go. Gotowe.
- **Nie zdążyłeś pobrać?** Discover → `bielik` → `speakeash/Bielik-4.5B-v3.0-Instruct-GGUF` (Q8_0, ~4,7 GB) — albo weź z **pendrive'a**, który krąży po sali.
- **Słaby laptop (≤8 GB RAM)?** Weź mniejszy `Bielik 1.5B` (~1,7 GB) — szybszy, choć skromniejszy.

Nadal nic? Otwórz Chrome i wejdź na **bielikchat.pl** — bez logowania.

AKT II

Co tracisz, ufając płatnym AI?

Ryzyko #1 — wyciek danych

Konsumenckie ChatGPT, Claude i Gemini używają Twoich rozmów do trenowania.

- Domyślnie opt-in. Opt-out istnieje, ale jest schowany w ustawieniach
- Plany enterprise i business mają zero-retention, ale doktoranta na nie nie stać
- Twoja praca doktorska, draft, niepublikowane wyniki — wszystko ląduje w systemach firmy w USA albo Chinach
- *Wciąż*, nawet po opt-out: dane przetwarzane na serwerach poza EU

Pytanie: kto z Was wrzucił już draft rozdziału do ChatGPT? *(ręce w górę)* Macie kopię tam.

Ryzyko #2 — GDPR przy danych respondentów

Wywiad z respondentem to dane osobowe. Punkt.

- Imię, miejsce, opinie polityczne, charakterystyka — wszystko to wchodzi pod RODO
- Wrzucenie do ChatGPT albo ElevenLabs STT oznacza transfer do data processora poza EU bez DPA
- Ryzyko prawne dotyczy nie tylko Ciebie, ale całej uczelni
- Komisje etyki i bioetyczne pytają wprost: gdzie będą przetwarzane dane?

Odpowiedź „w OpenAI” znaczy: nie zatwierdzimy.

Ryzyko #3 — AI Act + DORA (już obowiązuje)

- AI Act obowiązuje od 1 sierpnia 2024, pełna egzekucja od 2 sierpnia 2026
- Uczelnie są klasyfikowane — część zastosowań AI to systemy wysokiego ryzyka
- Wewnętrzne polityki AI uczelni dopiero powstają. Do tego czasu używasz GPT na własną odpowiedzialność
- Sankcje sięgają 7% globalnego obrotu, a dla uczelni — całego budżetu

W 2026 obraz tylko się zaostrzy. Suwerenne AI to inwestycja w przyszłą zgodność.

Ryzyko #4 — koszt

Pełen stack doktoranta to ChatGPT Plus, NotebookLM Pro, DeepL Pro, Grammarly Premium, ElevenLabs STT, Elicit i Notion AI.

W sumie 80-150 USD miesięcznie, czyli 400-700 zł.

Stypendium doktoranckie: 2700-3500 zł netto.

To ćwierć Twojej pensji co miesiąc. Plus VAT, plus kurs dolara.

Ryzyko #5 — halucynacje i fałszywe cytowania

Modele halucynują cytowania.

- Wymyślają autorów, tytuły, DOI i lata wydania
- W pracy naukowej kosztuje karierę — recenzent znajdzie nieistniejące źródło
- Realny przykład: amerykańscy prawnicy ukarani za cytowanie zmyślonego orzecznictwa wygenerowanego przez ChatGPT (Mata v. Avianca, 2023)

Antidotum: pracuj na realnym tekście, który **sam wklejasz**. Model poprawia, streszcza i tłumaczy to, co widzi — nie zmyśla źródeł, bo dostał je od Ciebie.

Dziś ćwiczymy dokładnie to.

Ryzyko #6 — „AI-style” i cenzura

Globalne modele piszą po polsku rozpoznawalnie. „Delve into”, „tapestry”, listy z em-dashami, charakterystyczne kalki językowe. Recenzenci coraz częściej rozpoznają ten styl — i mają wobec niego silne uprzedzenia. Detektory AI nie są wiarygodne, ale skutki dla Ciebie są realne.

Cenzura RLHF z kolei blokuje:

- Badania nad mową nienawiści, ekstremizmem, propagandą
- Prawnicze case studies z poważnymi przestępstwami
- Kwerendy historyczne (literatura okupacyjna, totalitarna)

Bielik trenowany na polskim korpusie pisze inaczej i nie odmawia analizy tekstu.

Akt II → III

Mapa: płatne → open-source

Mapa narzędzi: płatne → OSS + kompromisy (1/2)

Płatne	OSS odpowiednik	Główny kompromis
ChatGPT/Claude czat	Bielik 1.5/4.5/11B w LM Studio	jakość niższa na EN, brak agentów, wolniej
NotebookLM ★	LM Studio + Bielik ★	brak podcast-generatora, słabszy reranker
Elicit/scite wyszukiwarka	OpenAlex API + Zotero + lokalny RAG	musisz mieć PDF-y lokalnie, brak meta-grafu cytowań
Perplexity	SearxNG + Open WebUI + Bielik	wymaga hostingu, brak agentic browsing
DeepL Pro PL↔EN	Bielik 11B / NLLB-200	na egzotycznych językach DeepL wciąż lepszy
Grammarly	LanguageTool + Bielik	słabszy "tone detector" na EN

Mapa narzędzi: płatne → OSS + kompromisy (2/2)

Płatne	OSS odpowiednik	Główny kompromis
ElevenLabs STT transkrypcja	Whisper.cpp + Buzz GUI / WhisperX	wolniej (5-15 min na 1h audio), słabsza diaryzacja
Adobe Acrobat AI OCR	Marker / olmOCR / Tesseract	rękopisy trudne, brak GUI
Notion AI second brain	LM Studio: chat-per-projekt ★ / Obsidian + plugin	brak cloud sync (Syncthing/Git replace)
Gamma generator slajdów	Marp + Bielik do generowania MD	wymaga komfortu z Markdownem
DALL-E/Midjourney	Stable Diffusion XL / Flux	wymaga GPU 8+ GB VRAM
Zotero	<i>Już free i open-source!</i>	minimalne — używaj

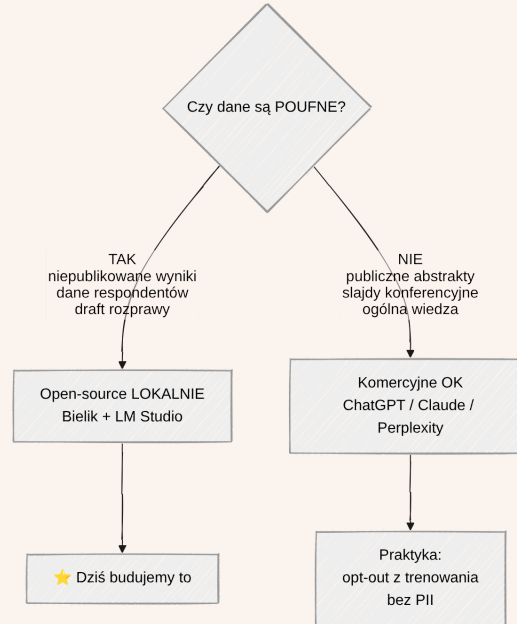
★ = używamy w warsztacie

80% funkcjonalności za 0 PLN

Kompromis: ~20-30% mniejsza jakość, czasem wolniej.

Zysk: pełna prywatność, koszt 0, GDPR-clean by design.

Kiedy hybryda jest OK?



Praktycznie:

- Burza mózgów nad tematem grantu? ChatGPT OK.
- Analiza wywiadu z respondentem? Tylko lokalnie.
- Polerka opublikowanego już abstraktu? Szara strefa, można w chmurze.
- Polerka niepublikowanego rozdziału? Tylko lokalnie.

AKT III

Hands-on: Bielik czyści Twoje dane

Jak rozmawiać z Bielikiem — kopiuje, wkleja, prosi

Cała dzisiejsza praktyka to jeden ruch, powtórzony kilka razy:

1. **Kopiujesz** kawałek tekstu — z pliku, maila, PDF-a, skanu po OCR.
2. **Wklejasz** go do czatu w LM Studio.
3. **Prosisz** Bielika, żeby coś z nim zrobić: poprawił, streścił, zanonimizował, przetłumaczył.

Żadnych plików do wgrywania, żadnych ustawień. Zwykła rozmowa po polsku.

Dobry prompt = konkretne, krótkie zadanie + tekst pod spodem.

Zamiast „pomóż mi z tym” → „popraw błędy OCR w poniższym tekście, zachowaj sens, nie dodawaj nic od siebie”.

Ćwiczymy na gotowym korpusie „**Poznański Czerwiec 1956**” — celowo brudnym (zepsute skany, surowe relacje, niespójne tabele). Te same prompty zadziałają potem na Twoich danych.

Krok 0 — załaduj Bielika do czatu

W LM Studio:

1. Lewy panel →  **Chat**
2. Na górze okna → „Select a model to load” → wybierz **Bielik 4.5B v3.0 Instruct** (na słabym laptopie: 1.5B).
Albo skrótem z Discover: „ Use in New Chat”.
3. Poczekaj 10-20 sekund, aż status zmieni się na zielony „Loaded ✓”.
4. W prawym panelu rzuć okiem na Prompt Template — powinno być ChatML (wykrywa się automatycznie z metadanych GGUF).
5. Nazwij chat (ikona ołówka przy chacie po lewej) — np. „Warsztat Bielik”.

W LM Studio każdy chat działa jak osobny pokój — trzymasz w nim model i całą historię rozmowy. Możesz mieć kilka chatów obok siebie, po jednym na zadanie albo projekt.

LM Studio

Update available

bielik-4.5b-v3.0-instruct

Chats

Neural Assistant Helps

relacja-janusz-bieganski_pol.txt
Plain text

Search chats...

New Folder

Warsztaty

Neural Assistant Help

Bielik 1.5: Szybkość i

Type to filter models...

Currently Loaded (1)
Memory Consumption: 270.67 MB / 30 GB

bielik-4.5b-v3.0-instruct Context: 8192 GGUF Eject

Your Models

				Recency	Size	Downloaded
Bielik 4.5B v3.0 Instruct	Q8_0	speakeash	4.8B	llama	GGUF	4.71 GB
Bielik 1.5B v3.0 Instruct	Q8_0	speakeash	1.6B	llama	GGUF	1.58 GB

☐ Manually choose model load parameters (or hold Alt)

Podsumuj mi ten wywiad

1. **Rozrywka i zaangażowanie młodzieży:** Młodzież (głównie studenci) uczestniczyła w działaniach przeciwko władzom PRL, takich jak podpalanie cystern z benzyną i obrzucanie budynków butelkami z paliwem. Starsi pracownicy opuszczali miejsca działań przed 2:00-3:00 nad ranem.

2. **Cel akcji:** Celem było zniechęcenie władz do represji wobec delegatów na zjazd ZSP (Związku Studentów Polskich), którzy mieli się spotkać w UB (Urząd Bezpieczeństwa). Młodzież chciała, aby delegaci nie pojawili się na spotkaniu lub żeby ich nie zastrzelono.

3. **Reakcja władz:** Władze odpowiedziały użyciem siły – ...gi i ostrzał ulic Poznania. Bądański wspomina o

Send a message to the model

Twój korpus: „Poznański Czerwiec 1956”

Tekst do ćwiczeń

-  `wycinki-prasowe/` — 3 wycinki z błędami OCR + skan
-  `relacje-swiadkow/` — 3 surowe transkrypcje wywiadów
-  `statystyki/ofiary-i-aresztowania.csv` — niespójna tabela
-  `opracowania/` — fragment PL + fragment EN (do tłumaczenia i porównania)

Audio

-  `nagrania/wspomnienia/` — 5 relacji świadków (.mp3)
-  `nagrania/przemowienie-cyrankiewicza-*.mp3` — archiwalne przemówienie

Już zamienione na tekst *(bez logowania)*

-  `nagrania/wspomnienia/relacja-janusz-bieganski_pol.txt`

Fotografie archiwalne

-  `fotografie/` — ~50 zdjęć z AIPN Poznań
 - demonstracje, walki uliczne
 - ofiary i ranni, procesy, zatrzymani

Źródła

- Wszystkie materiały z `ZRODLA.md` — licencje i cytowania

Stacja 1 — popraw zepsuty skan (OCR)

Otwórz `wycinki-prasowe/wycinek-01-ocr.txt` (Notatnik / TextEdit), zaznacz wszystko (Ctrl+A / ⌘A), skopiuj. Wklej do czatu, a **nad** wklejonym tekstem dopisz:

Popraw błędy OCR w poniższym tekście. Odtwórz poprawną polszczyznę z polskimi znakami (ą, ę, ł...), zachowaj oryginalny sens. Nie dodawaj nic od siebie.

To komunikat prasowy z 1956 r. „przepuszczony przez zły OCR” — `1` zamiast `l`, `0` zamiast `o`. Bielik odtwarza czytelny tekst w kilka sekund.

Realne zastosowanie: skany z archiwum, stare PDF-y, zdjęcia dokumentów.

Stacja 2 — rozmowa z dyktafonu → czysty tekst

Otwórz `relacje-swiadkow/relacja-01-surowa.txt`, skopiuj, wklej, a nad tekstem:

To surowa transkrypcja nagrania świadka. Dodaj interpunkcję, podziel na wypowiedzi (Świadek / Pytający), usuń potknięcia mowy („yyy”, „no więc”). Na końcu dodaj streszczenie w 5 zdaniach.

Zapis bez kropek, z „yyy” i wtrąceniami zmienia się w czytelny, podzielony tekst plus streszczenie.

Realne zastosowanie: wywiady, focus groupy, posiedzenia, zeznania.

Stacja 3 — anonimizacja (RODO)

Zostaw tę samą relację w czacie. Napisz pod spodem:

Zanonimizuj powyższy tekst zgodnie z RODO. Zamień imiona, nazwiska, adresy i inne dane osobowe na neutralne placeholdery: [ŚWIADEK_1], [ULICA], [MIASTO]. Zachowaj sens. Na końcu wypisz listę zmian.

Bielik znajduje dane osobowe (Tadeusz Kowalczyk, ul. Dąbrowskiego, żona Halina) i zastępuje je placeholderami — tekst gotowy do udostępnienia czy publikacji open data.

Dla prawników i nauk społecznych: to często warunek zgody komisji etyki.

Stacje bonus — wybierz pod siebie

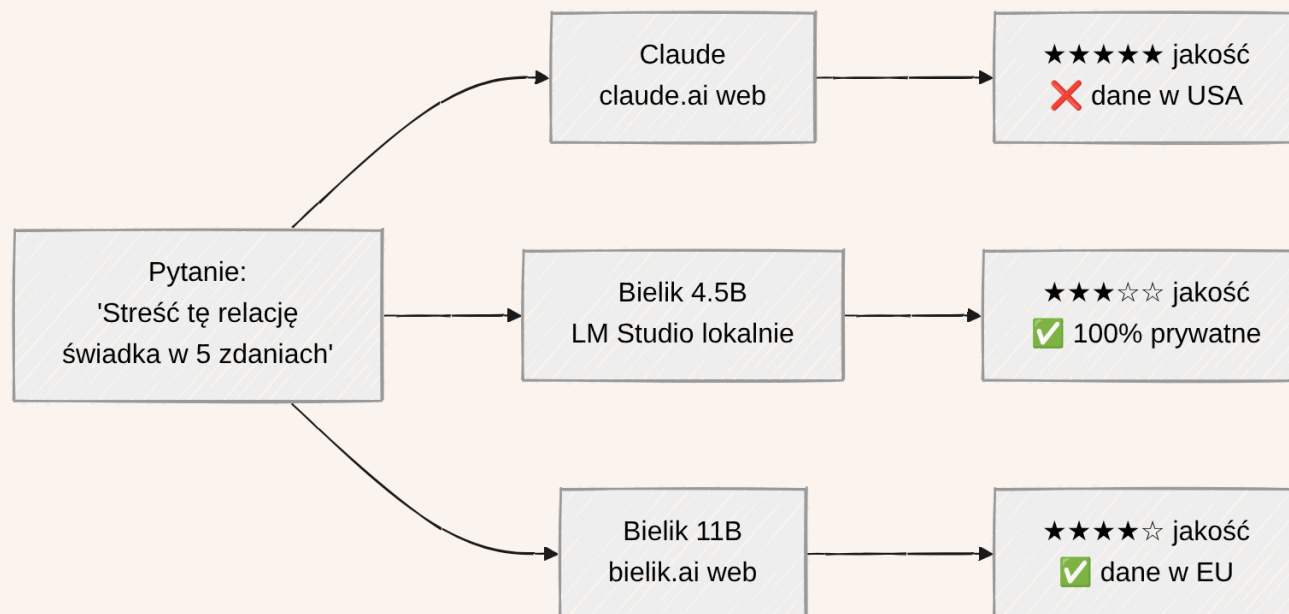
Skończyłeś szybciej? Pełne prompty masz w `prompty/` i w handout. Spróbuj:

Plik z korpusu	Poproś Bielika, żeby...
<code>statystyki/ofiary-i-aresztowania.csv</code>	uporządkował do jednej czystej tabeli: jeden separator, jeden format daty, bez duplikatów
<code>opracowania/study-en-excerpt.md</code>	przetłumaczył na polski ze spójną terminologią
oba pliki z <code>opracowania/</code>	porównał, ilu było zabitych — i wyjaśnił rozbieżność 57 vs 58

Ostatnia stacja pokazuje rzecz najważniejszą: AI **nie rozstrzyga** sprzecznych źródeł za Ciebie — pokazuje je obok siebie. Weryfikacja zostaje po Twojej stronie.

Demo: trzy poziomy Bielika

To samo pytanie zadane trzem różnym źródłom. Każde z innym kompromisem.



Pokazuję na żywo, split screen: po lewej LM Studio z Bielikiem 4.5B lokalnie, po prawej dwie zakładki przeglądarki — bielik.ai (11B) i claude.ai. Ten sam tekst (ta sama relacja z korpusu), to samo pytanie, trzy odpowiedzi obok siebie.

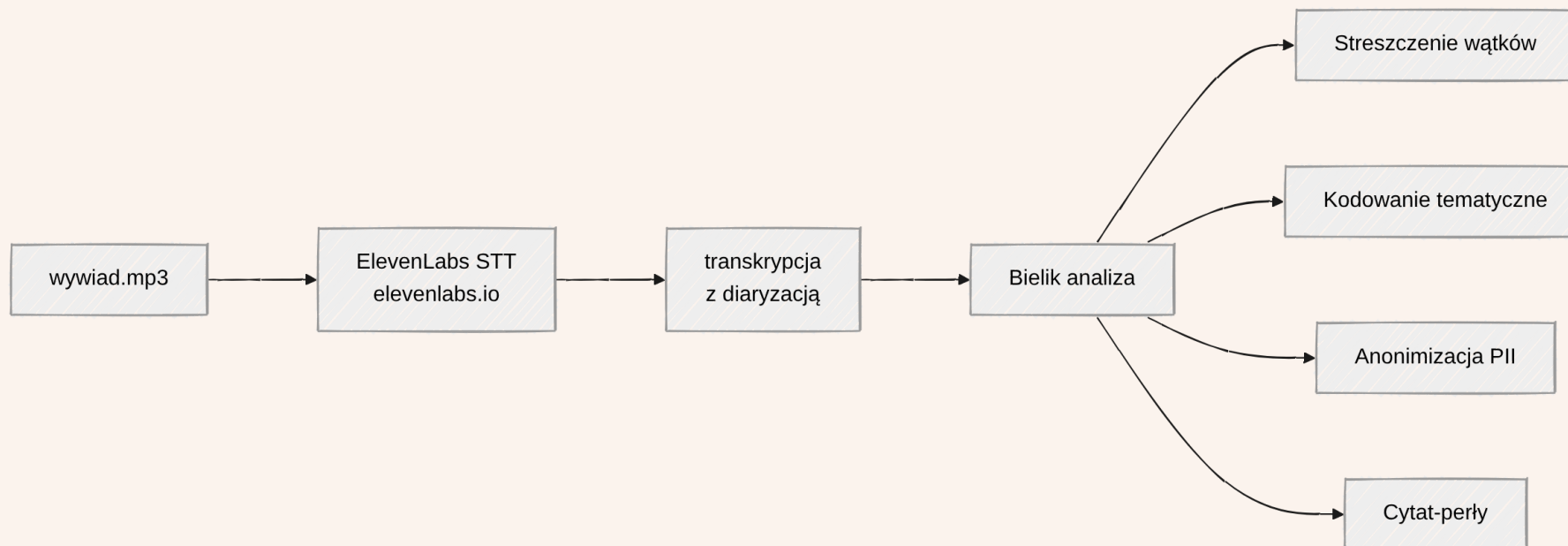
Provider-switch — wnioski

Provider	Jakość	Prywatność	Szybkość	Koszt	GDPR
Claude API	★★★★★	✗ USA	★★★★☆	\$\$\$	⚠
Bielik 4.5B lokalnie (domyślny)	★★★★☆	✓ na laptopie	★★★★☆☆	0	✓
Bielik 1.5B lokalnie (słaby sprzęt)	★★★☆☆	✓	★★☆☆☆☆	0	✓
Bielik 11B via bielik.ai	★★★★☆	✓ EU	★★★★☆	\$	✓

Wybór nie jest zero-jedynkowy. Zaczynasz od 4.5B (mieści się w 16 GB RAM i robi czyszczenie wiernie), wskazujesz na 11B przez bielik.ai gdy potrzebujesz najwyższej jakości, 1.5B trzymasz tylko na bardzo słaby laptop, ChatGPT — do danych publicznych.

Demo: transkrypcja wywiadu (ElevenLabs STT + Bielik)

Audio do tekstu przez przeglądarkę, tekst do analizy przez Bielika.



W Stacji 2 czyścieś **gotowy** tekst relacji. Tu pokazuję, skąd on się bierze: wchodzę na elevenlabs.io/pl/speech-to-text/polish, wrzucam nagranie z korpusu, dostaję transkrypt w kilkanaście sekund — i wklejam go do Bielika. *(Używamy ElevenLabs zamiast lokalnego Whispera — instalacja Whispera to kilkanaście minut, których dziś nie mamy.)*

Nie chcesz się logować? W materiałach jest gotowy plik `nagrania/rełacja-janusz-transkrypt.txt` — już zamieniony na tekst. Wklej go od razu do Bielika.

Open mic — z czym przyszedłeś?

Powiedz głośno:

- Jaka dyscyplina?
- Jakie dane chciałbyś przeanalizować?
- Co Cię boli w obecnym workflow?

Ja wtedy wskażę pozycję z tabeli płatne → OSS pod Twoją potrzebę, konkretną aplikację i prompt do testów, i powiem szczerze: czy zmieścimy POC w 90 minut, czy to weekend project.

Nie zmieścimy się ze wszystkimi. Kto nie zdąży — łapie mnie w przerwie albo mailem.

Materiały do zabrania

Starter-pack (link → QR na końcu):

-  `konfiguracja/lm-studio-config.md` — gotowe ustawienia chatu (*prompt template, długość kontekstu, warianty Bielika*)
-  `przykład-korpus/poznan-1956/` — gotowy korpus „brudnych” danych (skany, surowe relacje, niespójne tabele) + prawdziwe foto i audio
-  `prompty/` — gotowe szablony stacji (OCR, transkrypcja, anonimizacja, CSV, tłumaczenie, porównanie źródeł)
-  `przykład-korpus/poznan-1956/nagrania/` — prawdziwe nagrania (relacje świadków) — wrzuć do ElevenLabs STT i wklej transkrypt do Bielika
-  `instalacja/` — screenshoty Win/Mac/Linux + szybki start
-  **Handout 2-stronicowy** — 3 prompty rdzenia + tabela płatne → OSS + URL-e

Pełne slajdy (PDF) + ten warsztat na video → dostaniesz mailem w niedzielę.

Wracaj do Eskadry Bielika

Eskadra Bielika to program szkoleniowy SpeakLeash dla developerów i naukowców. Cztery rzeczy warto dopisać na po-warsztacie:

- Discord Eskadry, kanał `#naukowcy` — pilnuję go osobiście.
- `bielik.ai` — czat i API.
- HuggingFace SpeakLeash — repozytorium modeli.
- Open clinic za tydzień na Google Meet — pomagam z setupem u siebie.

Promocja warsztatu: tag `#BielikDlaDoktoranta` w mediach społecznościowych. Jeśli zbudujesz coś fajnego, daj znać —



Skanuj → Discord Eskadry
`discord.gg/eskadra-bielika`



Wszystko w jednym miejscu

<https://workshops.mmx3.pl/konferencja-krd-2026/>

 **Pobierz starter-pack ZIP:** <https://workshops.mmx3.pl/konferencja-krd-2026/starter-pack.zip>

Kontakt: max@mmx3.pl · LinkedIn [/in/maxmalecki](https://www.linkedin.com/in/maxmalecki)

Q&A

Zadaj pytanie na żywo — zeskanuj QR lub wejdź na Sli.do:

<https://app.sli.do/event/tf3xpZ9VrJtyam8zmGGBML>



Dziękuję!

Wszystko zostaje w laboratorium.

Nic nie wychodzi do BigTechów.

Backup #1 — co jeśli LM Studio nie startuje

Plan B — przeglądarka.

1. Otwórz Chrome / Firefox
2. Wejdź na: **bielikchat.pl**
3. Bielik czeka — wklejasz te same teksty i zadajesz te same pytania co reszta sali
4. Pracujesz przez przeglądarkę — *ta sama funkcjonalność, tylko nie na Twoim laptopie*

Bez logowania, bez instalacji. Działa natychmiast.

Backup #2 — Bielik na sprzęcie ze 4 GB RAM

Masz stary laptop? Da się.

- Bielik 1.5B Q4_K_M wciśnie się w 4 GB RAM, ledwie
- Generacja na CPU 8-tej generacji Intel: ~5 tokenów na sekundę. Wolno, ale działa
- Alternatywa: bielik.ai API za darmo (z rate-litem)
- Najprościej: **bielikchat.pl** — bez instalacji, bez logowania

Nie kupuj nowego laptopa pod AI. 8 GB RAM-u to komfort, 4 GB wystarczy.