

Lokalne systemy RAG i bazy wektorowe jako architektura suwerenności poznawczej w ujęciu Nitina Borwankara



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje koncepcję suwerenności poznawczej poprzez wdrażanie lokalnych modeli LLM i systemów RAG (Retrieval-Augmented Generation). Autor podkreśla, że przeniesienie infrastruktury AI z chmury na własny sprzęt (np. za pomocą Ollama) nie jest jedynie kwestią techniczną, lecz strategicznym wyborem w zakresie prywatności i kontroli nad danymi. Kluczem do sukcesu lokalnych RAG-u nie jest rozmiar modelu, lecz precyzyjne uziemienie go w kontrolowanym korpusie danych, rygorystyczny promet oraz niska temperatura generowania treści. Takie podejście zmienia rolę AI z „wszechwiedzącego proroka” w rzetelnego analityka dokumentów, minimalizując ryzyko halucynacji i zapewniając pełną audytowalność procesów poznawczych wewnątrz organizacji lub u użytkownika indywidualnego.

The article analyzes the concept of cognitive sovereignty through the implementation of local LLM models and RAG (Retrieval-Augmented Generation) systems. The author emphasizes that moving AI infrastructure from the cloud to one's own hardware (e.g., using Ollama) is not merely a technical issue, but a strategic choice regarding privacy and data control. The key to the success of local RAG is not the size of the model, but precisely grounding it in a controlled corpus of data, rigorous prompting, and low generation temperature. Such an approach changes the role of AI from an 'omniscient prophet' into a reliable document analyst, minimizing the risk of hallucinations and ensuring full auditability of cognitive processes within an organization or for an individual user.

Artykuł stanowi manifest przeciwko traktowaniu sztucznej inteligencji jako monolitycznego, nieprzejrzystego generatora odpowiedzi, proponując w zamian

przejście na architekturę systemową. Autor argumentuje, że prawdziwa wartość AI nie tkwi w elokwencji modeli językowych (LLM), lecz w ich ścisłym zintegrowaniu z lokalnymi bazami wektorowymi i mechanizmami RAG (Retrieval-Augmented Generation). Główna teza zakłada, że suwerenność poznawcza użytkownika zależy od przesunięcia kontroli z chmurowych ekosystemów na infrastrukturę lokalną, gdzie prymat prywatności, rygorystyczna wierność źródłom oraz standaryzacja zapytań (poprzez koncepcję VQL) przekształcają AI z „cyfrowej wyroczni” w audytowalne narzędzie wspomagania ludzkiego rozumu. Tekst dowodzi, że tylko systemy przejrzyste, wymuszające cytowania i podlegające rygorom logicznym mogą zapobiec abdykacji intelektualnej człowieka na rzecz halucynujących modeli. Tym samym architektura AI zostaje zdefiniowana nie jako wyścig po największy model, lecz jako budowa odpowiedzialnego ekosystemu pamięci, w którym technologia służy do odnajdywania sensu, a nie do jego pozorowania.

SŁOWA KLUCZOWE

lokalne systemy RAG

bazy wektorowe

suwerenność poznawcza

lokalne LLM

VQL Vector Query Language

architektura uziemienia

embeddingi

chunking

retrieval

halucynacje modeli AI

multimodalność

prywatność danych w AI

standardy zapytań wektorowych

audytowalność systemów AI

Ollama

Lokalne LLM jako fundament suwerenności poznawczej

Lokalne LLM. Prywatność, koszt i kontrola jako architektura, a nie dekoracja. W poprzednim fragmencie RAG został przedstawiony jako architektura uziemienia modelu językowego w źródłach. Pozostaje jednak pytanie: gdzie ten model ma działać? W chmurze komercyjnego dostawcy, gdzie wygoda jest wysoka, ale zależność równie duża? Czy lokalnie, na sprzęcie użytkownika lub organizacji, gdzie konfiguracja wymaga większej dyscypliny, lecz prywatność i kontrola nie są jedynie hasłem w prezentacji sprzedażowej? Borwankar wyraźnie skłania się ku drugiej ścieżce. W notatce źródłowej lokalne modele LLM opisano jako jeden z głównych tematów książki: sposób budowy osobistych systemów AI, które mogą działać z użyciem Ollama i modeli takich jak llama3.1:8b bez wysyłania prywatnych danych do zewnętrznych usług.

Ten wybór nie jest technologicznym kaprysem, lecz konsekwencją całej filozofii baz wektorowych jako infrastruktury pamięci. Jeżeli RAG ma pracować na prywatnych dokumentach, historiach rozmów, notatkach badawczych, archiwach organizacyjnych, dokumentacji prawnej lub danych wrażliwych, kwestia lokalności przestaje być pytaniem o wygodę. Staje się pytaniem o suwerenność poznawczą. Kto widzi dane? Kto przechowuje zapytania i analizuje wzorce użycia? Kto kontroluje logi, zmienia warunki dostępu lub ustala cenę w przyszłym kwartale? W świecie AI infrastruktura nie jest neutralnym przewodem – jest miejscem władzy.

Chmura oferuje potężne korzyści: dostęp do największych modeli, skalowalność, brak konieczności zarządzania sprzętem oraz szybkie wdrożenie. Jednak każda wygoda niesie ze sobą politykę ukrytą w regulaminie. Jeśli prywatne dokumenty muszą opuścić organizację, aby zostać przetworzone, powstaje nowa zależność. Może być ona akceptowalna i dobrze zabezpieczona, ale nie powinna pozostawać niewidzialna. Dojrzała architektura AI nie zaczyna się od zachwyty nad modelem, lecz od bilansu ryzyka: co wolno wysłać, czego nie wolno, co musi zostać lokalnie, a co można przetworzyć u dostawcy. Technologia pozbawiona takiego bilansu jest iluzorycznie zabezpieczona, a w praktyce komiczna. Lokalny LLM odpowiada na tę potrzebę, ale nie czyni tego bez kosztów. Praca lokalna wymaga odpowiedniego sprzętu, pamięci, miejsca na modele oraz umiejętności ich instalacji i testowania. Mniejsze modele lokalne mogą być mniej sprawne w rozumowaniu niż największe rozwiązania chmurowe; mogą gorzej radzić sobie z długim kontekstem, językiem specjalistycznym czy subtelną argumentacją. Nie wolno więc czynić z lokalności bożka. Lokalny system z kiepskim modelem, słabą bazą, błędnym chunkingiem i luźnym promptem będzie tylko prywatną maszyną do produkcji prywatnych błędów. Prywatnych, owszem, ale wciąż błędów.

Siła lokalnego RAG-u leży gdzie indziej. Nie w tym, że lokalny model jest zawsze „mądrzejszy”, lecz w możliwości ścisłego osadzenia go w kontrolowanym korpusie. Jeśli retrieval jest sprawny, baza dobrze przygotowana, dokumenty poprawnie pocięte, a prompt rygorystyczny, model nie musi posiadać encyklopedycznej wiedzy. Ma wykonać konkretne zadanie: przeczytać dostarczone fragmenty, uporządkować je i odpowiedzieć na pytanie, wskazując źródła bez dopisywania treści spoza kontekstu. W takim układzie lokalny model przestaje być cyfrowym prorokiem, a staje się pracownikiem kancelarii archiwalnej: nie musi znać całego świata, ale ma rzetelnie obsłużyć akta, które otrzymał na biurko. Lokalny RAG odwraca dominującą wyobraźnię o AI. Nie chodzi o największy model, który „wie wszystko”, lecz o wystarczająco sprawny model pracujący na właściwych danych. To zasadnicza różnica.

Wielkie modele językowe imponują rozległością kompetencji, ale w zadaniach instytucjonalnych, prawnych czy naukowych najważniejsza jest trafność względem korpusu. Przy zapytaniach o własne dokumenty, uchwały i analizy ogólna wiedza świata staje się drugorzędna wobec zdolności

odnalezienia właściwego fragmentu i wiernego jego przytoczenia. Borwankar zaleca w takich zastosowaniach bardzo niską temperaturę modelu, na przykład 0.1. To parametr techniczny, ale jego sens jest niemal moralny: ograniczenie pokusy twórczej improwizacji tam, gdzie niezbędna jest wierność. Wysoka temperatura służy brainstormingowi i stylizacji; jednak gdy model analizuje dokumentację naukową lub prywatne notatki, nadmierna kreatywność prowadzi do błędów merytorycznych, które w rzeczywistości okazują się kosztowne.

Drugim narzędziem dyscyplinowania LLM jest prompt. Instrukcje muszą jasno definiować rolę modelu, zakazywać korzystania z wiedzy zewnętrznej i wymuszać cytowanie źródeł dla każdego twierdzenia. Dobry prompt w systemie RAG nie jest uprzejmą prośbą, lecz regulaminem pracy. Ma ograniczać model, a nie stymulować jego popisy. W generatywnej AI jakość często zależy od tego, czego skutecznie zabronimy. Wolność bez procedury bywa piękna w poezji, lecz w analizie dokumentów kończy się halucynacją w garniturze. Wymóg cytowania pełni tu funkcję szczególną: każda odpowiedź ma być sprawdzalna. Model powinien wskazywać konkretny fragment źródła zamiast sugerować prawdopodobieństwo. W systemach lokalnych daje to dodatkową zaletę – użytkownik może samodzielnie kontrolować cały łańcuch od dokumentu do odpowiedzi.

Lokalna architektura jako system komponentów i kontroli

Użytkownik wie, jakie pliki znajdują się w bazie, jak zostały przetworzone, które fragmenty pobrano i co model z nimi zrobił. Taka przejrzystość jest trudniejsza do osiągnięcia w systemach, gdzie retrieval, ranking, model i logika bezpieczeństwa pozostają ukryte za komercyjnym interfejsem. Nie wszystko, co wygodne, jest audytowalne, a nie wszystko, co audytowalne, musi mieć ładny przycisk.

Lokalność wiąże się jednak z kosztami operacyjnymi. Korzystanie z zewnętrznych modeli przez API bywa wygodne, lecz przy dużej liczbie zapytań, masowym przetwarzaniu dokumentów czy eksperymentach badawczych wydatki mogą szybko wzrosnąć. Lokalny model, raz pobrany i uruchomiony, nie generuje kosztu za każde zapytanie w tradycyjnym sensie. Oczywiście istnieją koszty sprzętu, energii, utrzymania oraz czasu specjalistów, ale ekonomia tego rozwiązania jest inna. Dla indywidualnych badaczy, małych zespołów, fundacji, redakcji czy organizacji, które nie chcą płacić za każde potknięcie promptu, lokalny stos technologiczny może być atrakcyjny – zwłaszcza gdy eksperymentowanie jest częścią pracy, a nie incydentalną fanaberią.

Nie należy jednak mylić braku opłaty za zapytanie z całkowitym brakiem kosztów. Lokalny system wymaga utrzymania: modele trzeba aktualizować, bazy indeksować, a dokumenty wersjonować.

Sprzęt się starzeje, pamięć zapelnia, a pipeline psuje się przy zmianie formatu pliku lub biblioteki. Kto obiecuje lokalną AI jako rozwiązanie „bez kosztów”, sprzedaje bajkę z kartą graficzną w roli różdżki. Uczciwa teza brzmi inaczej: lokalność przesuwa koszty z opłat operacyjnych i zależności od dostawcy na kompetencje, sprzęt i utrzymanie. Dla wielu zastosowań to przesunięcie jest korzystne, choć musi być świadomie rozumiane.

W koncepcji Borwankara lokalne LLM są naturalnym przedłużeniem lokalnych baz wektorowych. Jeśli cały system wiedzy ma działać na prywatnym komputerze lub w kontrolowanej infrastrukturze, to baza SQLite, PostgreSQL z pgvector, FAISS, SentenceTransformers i Ollama tworzą spójny ekosystem. Każdy element można zastąpić innym, lecz filozofia pozostaje ta sama: minimalizować zbędne zależności, maksymalizować kontrolę nad danymi i budować systemy zrozumiałe oraz odtwarzalne. To jest duch open-source: nie tylko darmowość, lecz możliwość wglądu w mechanizm. W epoce AI wgląd w mechanizm jest nową formą alfabetyzacji.

Zrozumienie tej architektury wymaga rozróżnienia trzech typów modeli językowych: encoder-only, decoder-only oraz encoder-decoder. Modele encoder-only, takie jak BERT, specjalizują się w rozumieniu i klasyfikacji tekstu; można je traktować jako mechanizmy głębokiego odczytu reprezentacji. Modele decoder-only, do których należą rodziny GPT i Llama, są zoptymalizowane pod generowanie kolejnych tokenów, a więc tworzenie płynnego tekstu. Z kolei modele encoder-decoder, takie jak T5 lub BART, łączą obie funkcje, sprawdzając się w zadaniach transformacji: tłumaczeniu, streszczaniu czy parafrazowaniu. Ten podział dowodzi, że „model językowy” nie jest monolitem, lecz rodziną narzędzi o różnych talentach i wadach.

W systemach RAG typowo stosuje się osobny model embeddings do wyszukiwania i osobny LLM do generowania odpowiedzi. Jest to kluczowe, gdyż duży model generatywny nie musi być najskuteczniejszym ani najtańszym narzędziem do tworzenia wektorów. Współcześnie rzadko wykorzystuje się potężne LLM do samego generowania embeddings; częściej stawia się na mniejsze, wyspecjalizowane modele, takie jak all-MiniLM-L6-v2. To klasyczny podział pracy: jeden model zamienia tekst na przestrzeń semantyczną, inny odpowiada językiem.

Próba realizacji wszystkich zadań za pomocą jednego narzędzia może być koncepcyjnie wygodna, lecz rzadko efektywna. Tu pojawia się fundamentalna zasada projektowa: system AI jest architekturą komponentów, nie monolitem. Model embeddingowy, baza wektorowa, indeks, wyszukiwarka hybrydowa, reranker, LLM, prompt, orkiestrator i interfejs użytkownika pełnią odrębne funkcje.

Lokalne LLM jako asystent rozumu, a nie wyrocznia

Zmiana jednego elementu może poprawić lub pogorszyć całość; dojrzałość inżynierska polega więc na rozumieniu przepływu informacji między komponentami. Kto sugeruje jedynie „użycie lepszego modelu”, często przypomina lekarza, który na każdą chorobę przepisuje większy młotek – rozwiązanie spektakularne, lecz diagnostycznie ubogie. Lokalne modele LLM są kluczowe również przy pracy z dokumentacją wielojęzyczną i domenową. W polskich realiach dotyczy to zwłaszcza tekstów prawnych, administracyjnych, publicystycznych czy naukowych, które cechują się specyficzną składnią i terminologią. Model ogólny może władać językiem polskim, ale niekoniecznie rozumie niszowy język uchwał, petycji, przepisów, argumentacji konstytucyjnej czy orzecznictwa. Tutaj pomocny okazuje się retrieval: nawet jeśli model nie posiada silnej wiedzy domenowej w swoich parametrach, może otrzymać właściwe fragmenty w kontekście. RAG nie czyni go ekspertem z natury, ale wymusza pracę na materiale eksperckim.

Największym wyzwaniem pozostaje jednak wiarygodność odpowiedzi. Lokalny model zapewnia większą prywatność, lecz nadal może halucynować. Może błędnie złożyć fragmenty, pominąć ograniczenie, zbyt mocno uogólnić, przeoczyć sprzeczności między źródłami lub nadać zbyt pewny ton niepewnej odpowiedzi. Dlatego system należy projektować tak, aby model miał obowiązek rozdzielać fakty od interpretacji i braków w danych. Powinien umieć stwierdzić: „dostarczone źródła wskazują X”, „nie znalazłem podstawy dla Y”, „to jest interpretacja, a nie cytowany fakt” lub „źródła są sprzeczne”. Taki sposób odpowiedzi jest mniej efektowny niż pewny monolog, ale znacznie bardziej profesjonalny. W nauce, prawie i analizie instytucjonalnej poprawnie oznaczona niepewność jest cenniejsza niż udawana pewność. Lokalny RAG powinien być zatem narzędziem wspomagania rozumu, a nie jego protezą. Użytkownik nie powinien traktować modelu jak wyroczni, lecz jak asystenta, który pomaga przeszukać korpus, wskazać fragmenty, uporządkować materiał i skrócić drogę do konkretnego dokumentu.

Ostateczna ocena musi pozostać po stronie człowieka. W przeciwnym razie przechodzimy od problemu halucynacji modelu do problemu abdykacji człowieka. Abdykacja poznawcza jest znacznie groźniejsza, bo model przynajmniej nie udaje obywatela. Lokalne LLM mogą odegrać szczególną rolę w edukacji i badaniach: pozwalają eksperymentować bez ciągłego liczenia tokenów, budować własne korpusy, testować strategie chunkingu, porównywać modele embeddings czy sprawdzać progi podobieństwa. Dla programisty, badacza lub analityka to ogromna wartość – nie uczy się wtedy jedynie obsługi gotowego narzędzia, lecz rozumie, co faktycznie dzieje się między dokumentem a odpowiedzią. To różnica między kierowcą aplikacji a mechanikiem silnika. Obaj mogą dojechać do celu, ale tylko jeden rozumie, dlaczego samochód nagle zaczął brzmieć jak blender pełen śrubek.

Lokalne systemy jako gwarant różnorodności i podmiotowości

Równocześnie lokalność może wspierać pluralizm technologiczny. Rynek AI wykazuje silną tendencję do koncentracji: dominują największe modele, centra danych, dostawcy i budżety. Lokalne modele, open-source, małe bazy wektorowe oraz prywatne systemy RAG stanowią przeciwwagę dla tej hegemonii. Choć nie zastąpią one w pełni infrastruktury gigantów, pozwalają zachować przestrzeń do eksperymentu, niezależności i dostosowania narzędzi do konkretnych potrzeb. W świecie, w którym każdy dąży do bycia platformą, własny stos technologiczny jest czasem jak własny warsztat drukarski w epoce wielkich wydawnictw.

Ta perspektywa ma również wymiar kulturowy. Jeśli rozwój AI zostanie powierzony wyłącznie kilku globalnym dostawcom, lokalne języki, niszowe domeny, mniejsze instytucje czy specyficzne tradycje argumentacji będą zależne od tego, czy ich obsługa okaże się rentowna. Lokalne systemy umożliwiają budowanie pamięci cyfrowej dopasowanej do konkretnych wspólnot wiedzy. Pozwalają stworzyć RAG dla archiwum fundacji, dokumentów samorządowych, historii lokalnej, języka prawnego, korpusu religioznawczego czy badań nad niszową dziedzina.

Nie jest to kwestia marginesu, lecz warunek różnorodności poznawczej. W tym miejscu pojawia się wątek modeli multimodalnych – naturalny kierunek rozwoju hybrydowych systemów wiedzy. Dokumenty nie składają się przecież wyłącznie z tekstu; publikacje naukowe zawierają wykresy i tabele, a organizacje operują na skanach, prezentacjach czy plikach dźwiękowych. Modele multimodalne mogą tworzyć wspólne reprezentacje dla różnych typów danych, pozwalając systemowi odnajdywać znaczenie nie tylko w zdaniach, ale i w strukturach wizualnych. Jest to kluczowe zwłaszcza w nauce, gdzie istotne informacje zawarte w wykresach pozostają niewidoczne dla klasycznego parsera tekstu.

Multimodalność zwiększa jednak stopień złożoności. O ile tekst można pociąć na chunki (choć nie zawsze jest to proste), o tyle obraz wymaga innej reprezentacji, tabela – zachowania relacji między komórkami, a wykres – rozpoznania osi i legendy. Wspólna przestrzeń wektorowa dla wielu modalności jest obiecująca, lecz ryzykowna: łatwo pomylić wizualne podobieństwo z podobieństwem znaczenia.

W multimodalności diabeł nie tylko zrobił doktorat. On prowadzi interdyscyplinarne centrum badawcze. Lokalne systemy multimodalne będą więc wymagały jeszcze większej dyscypliny audytu. Gdy system opiera się na tekście, weryfikacja cytatu jest stosunkowo prosta; gdy bazuje na obrazie lub tabeli, sprawdzalność staje się wyzwaniem. Konieczne będzie opracowanie mechanizmów

precyzyjnego wskazywania regionów obrazu czy ramek czasowych nagrania. W przeciwnym razie multimodalny RAG może generować błędy jeszcze bardziej przekonujące i trudniejsze do wykrycia. Im więcej modalności, tym większy obowiązek przejrzystości.

Mimo tych ryzyk, kierunek lokalnych systemów pozostaje fundamentalny. Daje szansę na budowę narzędzi prywatnych, audytowalnych i edukacyjnych. Zamiast czekać na zainteresowanie globalnej platformy naszą niszową potrzebą, możemy zbudować własny korpus, pipeline i asystenta.

Nie zawsze będzie to łatwe, ale kompetencja rzadko rodzi się z wygody. Łatwość rodzi użytkowników. Trudność, jeśli dobrze prowadzona, rodzi podmiotowość. W tym kontekście lokalny LLM nie jest końcem opowieści, lecz etapem prowadzącym do pytania o standardy. Jeśli coraz więcej organizacji zacznie budować własne systemy wektorowe, pojawi się problem fragmentacji. Różne API, nazwy operacji i formaty zapytań sprawią, że programiści będą pisać kod zależny od konkretnego dostawcy, a badacze stracą możliwość powtarzalności eksperymentów. Rynek zacznie przypominać wieżę Babel, w której zamiast języków naturalnych będziemy operować dialektami SDK.

VQL jako niezbędna gramatyka dla dojrzałości baz wektorowych

W tym kontekście VQL (Vector Query Language) jawi się jako logiczny następny krok. Skoro wektory stały się nowym typem danych, similarity search nową operacją podstawową, metryki odległości częścią logiki zapytań, a RAG wymaga powtarzalnych procedur retrieval – potrzebujemy języka, który pozwoli opisywać te procesy w sposób jednolity. Nie po to, by zamknąć innowację w biurokratycznej klatce, lecz by nadać jej wspólną gramatykę. Bez niej każdy mówi po swojemu, a wtedy zwykle wygrywa ten, kto dysponuje największym megafonem, a nie ten, kto posiada najlepszą składnię. VQL w ujęciu Borwankara pozostaje propozycją teoretyczną i eksperymentalną, opartą na SQL, a nie gotowym, oficjalnym standardem.

Nie należy zatem przedstawiać go jako języka produkcyjnego, który już dominuje na rynku. Jego znaczenie polega raczej na wyznaczeniu kierunku: bazy wektorowe dojrzeją w momencie, gdy przestaną być zbiorem osobnych interfejsów i zaczną tworzyć wspólny język zapytań. Podobnie jak SQL stał się fundamentem pracy z danymi relacyjnymi, pozwalając deklaratywnie określać potrzeby użytkownika względem bazy, tak VQL miałby umożliwić deklaratywne definiowanie oczekiwań wobec przestrzeni wektorowej.

Dążenie do standaryzacji nie kłóci się z lokalnością. Przeciwnie: lokalne systemy potrzebują standardów bardziej niż wielcy dostawcy. Gigant może narzucić własny ekosystem i zmusić użytkowników do adaptacji; mała organizacja natomiast potrzebuje przenośności, czytelności, łatwej migracji i niezależności od jednego narzędzia. Standard staje się narzędziem wolności wtedy, gdy chroni przed zamknięciem w cudzym dialekcie. Oczywiście może on stać się biurokratycznym betonem, jeśli powstanie zbyt wcześnie lub zostanie przejęty przez interesy największych graczy. Dlatego VQL powinien wyrastać z praktyki open-source i rzeczywistych przypadków użycia, a nie z komitetowej ambicji pisania konstytucji dla technologii, która jeszcze uczy się chodzić.

Czas zatem przyjrzeć się samej koncepcji VQL: jego motywacji, modelowi danych, składni, operacjom podobieństwa, wyszukiwaniu hybrydowemu, range search, batch operations, funkcjom wektorowym oraz agregacjom i politycznemu znaczeniu standaryzacji. VQL jest bowiem czymś więcej niż pomysłem na wygodniejszy zapis zapytań. To próba nazwania nowej epoki baz danych, w której pytanie nie brzmi już tylko „gdzie jest rekord?”, lecz także: „co jest znaczeniowo blisko, według jakiej metryki, z jakim progiem zaufania i pod czyją kontrolą?”.

Vector Query Language. Dlaczego przestrzeń wektorowa potrzebuje własnej gramatyki?

Po analizie wektorów, embeddings, FAISS, lokalnych baz, RAG-u i modeli LLM dochodzimy do pytania, które z pozoru dotyczy jedynie składni, a w rzeczywistości uderza w sedno dojrzałości całej infrastruktury AI: czy bazy wektorowe potrzebują własnego języka zapytań? Borwankar odpowiada na to pytanie poprzez koncepcję Vector Query Language. W notatce źródłowej wyraźnie podkreślono, że nie jest to obecnie w pełni zaimplementowany standard, lecz teoretyczny, eksperymentalny i inspirowany SQL prototyp, pomyślany jako punkt wyjścia do dyskusji branżowej i potencjalnej normalizacji.

VQL jako deklaratywny standard zapytań wektorowych

VQL nie jest jeszcze kodeksem cywilnym baz wektorowych. Przypomina raczej projekt konstytucji dla państwa, które dopiero odkryło, że składa się z wielu księstw – każde z własną walutą, miarą długości i urzędnikiem od podobieństwa kosinusowego. Dzisiejszy rynek baz wektorowych jest rozdrobniony: systemy oferują różne API, sposoby zapisu zapytań, nazewnictwo operacji czy metody filtrowania, a także odmienne domyślne metryki, strategie indeksowania i poziomy integracji z metadanymi. Dla początkującego użytkownika jest to chaos edukacyjny. Dla programisty – ryzyko vendor lock-in. Dla DevOps i SRE – konieczność utrzymywania osobnych skryptów i procedur.

Badacz mierzy się tu z problemem powtarzalności eksperymentu, a organizacja z trudnością migracji między rozwiązaniami. W praktyce oznacza to, że warstwa retrieval, która powinna być fundamentem systemów RAG, często bywa zaszyta w kodzie aplikacji tak głęboko, że staje się niewidzialna, nieprzenośna i trudna do audytu. Borwankar dostrzega tu analogię do roli, jaką SQL odegrał dla baz relacyjnych. SQL nie zlikwidował różnic między silnikami bazodanowymi, ale ustanowił wspólną gramatykę pracy z tabelami, rekordami, filtrami i agregacjami. Pozwolił użytkownikowi deklaruować, czego chce od danych, zamiast opisywać krok po kroku, jak silnik ma to wykonać. VQL ma spełnić podobną funkcję wobec danych wektorowych. Nie chodzi o zastąpienie SQL-a, lecz o rozszerzenie jego ducha na świat podobieństwa, metryk odległości, przestrzeni semantycznych, operacji top-k, progów trafności i wyszukiwania hybrydowego.

Istotą VQL jest zatem deklaratywność. Użytkownik nie powinien za każdym razem pisać procedury przeszukiwania zależnej od konkretnego dostawcy. Powinien móc określić: wybierz przestrzeń z danej kolekcji i tabeli, wykonaj operację wektorową przy użyciu wskazanej metryki, zastosuj filtry metadanych i ogranicz wynik do określonej liczby rezultatów lub odrzuć te poniżej progu zaufania. W notatce źródłowej podstawowy szablon zapytania VQL przedstawiono jako połączenie klasycznych klauzul `SELECT`, `FROM`, `WHERE` i `LIMIT` z nowymi elementami `VECTOR_OPERATION` oraz `USING METRIC`. Ten krótki zapis niesie głęboką zmianę: podobieństwo staje się pełnoprawną operacją języka danych. Klasyczny SQL pytał o zgodność warunków – czy rekord spełnia kryterium, czy wartość pola mieści się w przedziale lub czy kategoria jest równa wskazanemu tekstowi. VQL pyta inaczej: co jest blisko? Jak bardzo blisko? Według jakiej metryki i w jakiej przestrzeni? Z jakim filtrem, limitem i progiem?

To przejście od logiki identyczności do logiki podobieństwa. Nie oznacza ono jednak porzucenia twardych filtrów; przeciwnie – najskuteczniejsze zapytania wektorowe łączą miękką semantykę z twardą selekcją metadanych. Właśnie dlatego VQL łączy `WHERE` z operacjami wektorowymi, by baza rozumiała zarówno sens, jak i rubrykę. Ostatecznie nawet najbardziej poetycki embedding powinien czasem uszanować datę faktury. Model danych VQL obejmuje kolekcje, tabele, wektory, embeddings, pola, metadane oraz metryki odległości. Kolekcja jest najwyższym poziomem organizacji (odpowiednikiem bazy), tabela przechowuje ustrukturyzowany zbiór wektorów z metadanymi, a sam wektor jest tablicą liczbową o stałej liczbie wymiarów. Embedding to szczególny rodzaj wektora reprezentujący zakodowaną jednostkę znaczenia. Pole oznacza konkretną wartość wymiaru w przestrzeni, natomiast metadane to dane opisowe, które nie uczestniczą bezpośrednio w obliczeniach. Metryka odległości określa sposób mierzenia dystansu.

Taki model pokazuje, że VQL nie traktuje wektora jako przypadkowego BLOB-a, lecz jako natywny byt zapytaniowy. To ujęcie jest kluczowe. Przez długi czas dane nieustrukturyzowane były w bazach

obiektami drugiej kategorii: plikami czy binarnymi paczkami, z którymi system nie potrafił zrobić nic poza ich przechowaniem. Wektor zmienia tę sytuację – staje się obiektem, który można porównywać, dodawać, normalizować i agregować. Jeśli baza danych ma naprawdę wspierać AI, nie może traktować embeddings jako egzotycznych załączników.

Musi dać im status obywateli pierwszej klasy. VQL jest próbą napisania dla nich prawa miejskiego. Pierwszą główną operacją języka jest similarity search, czyli wyszukiwanie podobieństw. W klasycznym zastosowaniu system znajduje wektory najbliższe zapytaniu. Przykład z notatki źródłowej – wyszukiwanie produktów podobnych przy użyciu metryki euklidesowej i filtrów kategorii oraz ceny – dobrze obrazuje siłę VQL. Użytkownik nie szuka „dowolnych” podobnych produktów, lecz takich, które spełniają konkretne warunki biznesowe: cenę, kategorię czy dostępność w danym kraju. Samo podobieństwo semantyczne byłoby zbyt szerokie, a same filtry – zbyt sztywne. W praktyce similarity search stanowi fundament systemów rekomendacyjnych, wyszukiwarek dokumentów, wykrywania duplikatów, analizy kodu oraz architektury RAG.

VQL jako standard transparentnego i precyzyjnego wyszukiwania

VQL pozwoliłby zapisać tę operację w sposób czytelny i przenośny, co jest kluczowe z punktu widzenia audytu. Deklaratywny zapis zapytania umożliwia jego analizę, wersjonowanie, porównywanie oraz dokumentowanie. Gdy logika retrieval jest rozproszona w kodzie aplikacji, staje się znacznie trudniejsza do kontroli. W systemach generujących odpowiedzi na podstawie źródeł proces ten jest zbyt istotny, by ukrywać go jak wstydlivy skrypt w piwnicy projektu.

Drugą kluczową operacją jest thresholding, czyli zastosowanie progu odcięcia. Klauzula THRESHOLD zmusza system do odrzucenia wyników poniżej określonego poziomu zaufania lub podobieństwa, co ma fundamentalne znaczenie w architekturze RAG. Jeśli system pobierze fragmenty jedynie dlatego, że musi coś zwrócić – nawet przy niskim stopniu podobieństwa – model językowy może budować odpowiedź na kruchym materiale. Próg pełni tu rolę zasady minimalnej relewantności: lepiej nie odpowiedzieć lub zrobić to ostrożnie, niż opierać się na przypadkowych sąsiadach semantycznych. W epoce generatywnej elokwencji taka odmowa jest aktem higieny poznawczej.

Trzecią operacją jest hybrid search, czyli wyszukiwanie hybrydowe. VQL miałby natywnie obsługiwać łączenie wektorów ze standardowym wyszukiwaniem tekstowym (np. BM25). Przykład wag – wektor z wartością 0.7 i tekst z wartością 0.3, sumujące się do 1.0 – obrazuje istotę tej

koncepcji. Wyszukiwanie hybrydowe nie jest półśrodkiem dla niezdecydowanych, lecz uznaniem, że znaczenie posiada różne tryby obecności. Czasem kluczowy jest sens ogólny, czasem nazwa własna, dokładny symbol, pojęcie prawne, którego nie wolno zastąpić synonimem, lub fraza techniczna wymagająca literalnego dopasowania. Hybryda łączy elastyczność wektorów z precyzją słowa.

W tym miejscu pojawia się jeden z najważniejszych wymiarów projektowania systemów wiedzy: normalizacja wyników. Wyniki BM25 i podobieństwa wektorowego operują na różnych skalach; nie można ich naiwnie sumować, tak jak nie dodaje się temperatury do długości stołu, by wyliczyć średnią „komfortu mebla”. Wymagają one przeskalowania i normalizacji przed nadaniem wag. VQL jako język mógłby wymuszać świadome opisywanie tych operacji, dzięki czemu wyszukiwanie hybrydowe przestałoby być zbiorem ręcznych trików w kodzie, a stałoby się przejrzystą częścią zapytania. Dla profesjonalnych systemów to różnica między rzemiosłem a alchemią.

Kolejna operacja, range search, stanowi odejście od klasycznego modelu top-k. Podczas gdy top-k zwraca określoną liczbę najlepszych wyników niezależnie od ich jakości, range search dostarcza wszystkie wektory mieszczące się w zadanym promieniu lub poniżej progu odległości. Różnica jest subtelna, lecz zasadnicza. Top-k odpowiada na pytanie: „pokaż mi pięć najbliższych obiektów”, natomiast range search pyta: „pokaż mi wszystko, co spełnia kryterium bliskości”. W zastosowaniach wysokiego ryzyka to drugie podejście jest uczciwsze, gdyż nie udaje, że zawsze istnieje satysfakcjonująca piątka wyników. Czasem właściwa odpowiedź brzmi: w tej przestrzeni nie ma wystarczająco bliskich sąsiadów.

Operacje wsadowe (batch operations) rozwiązują z kolei problem efektywności przy wielu zapytaniach. BATCH SIMILARITY SEARCH pozwala przesłać do bazy wiele wektorów jednocześnie, co redukuje opóźnienia sieciowe i umożliwia zwrot pogrupowanych wyników w oryginalnej kolejności zapytań.

VQL jako standard analityki semantycznej i infrastruktura wolności

To brzmi technicznie, ale ma istotne znaczenie praktyczne. W realnych systemach rzadko pytamy o jeden wektor; często chcemy porównać setki dokumentów, wykryć duplikaty w korpusie, przypisać podobne sprawy do wielu nowych przypadków, przygotować masową analizę rekomendacji lub zasilić pipeline RAG dla wielu zapytań. Operacje wsadowe są językowym uznaniem faktu, że AI pracuje coraz częściej w rytmie produkcyjnym, nie demonstracyjnym.

VQL obejmowałby także funkcje wektorowe. W notatce wymieniono pobieranie wymiarowości wektora, obliczanie dystansu między dwoma wektorami, konkatencję, iloczyn skalarny, dodawanie wektorów, skalowanie przez mnożenie oraz agregacje, takie jak średnia wektorów. To interesujący fragment, gdyż pokazuje, że język zapytań nie powinien ograniczać się do wyszukiwania, lecz umożliwiać podstawową matematykę na reprezentacjach semantycznych. Jeśli wektor jest natywnym typem danych, operacje na nim powinny być dostępne w silniku bazy, a nie zawsze wypychane do zewnętrznego kodu aplikacji. Arytmetyka wektorowa ma swoje symboliczne źródło w przykładach Word2Vec, gdzie relacje semantyczne modelowano przez dodawanie i odejmowanie wektorów. Notatka wspomina o operacjach takich jak dodawanie `influence_vector` do `vector_data` lub skalowanie wektora przez wartość `scalar_value`. Należy tu jednak zachować ostrożność.

Arytmetyka wektorowa nie jest magiczną logiką pojęć i nie każda operacja matematyczna na embeddingach ma sens interpretacyjny. Niemniej w zadaniach takich jak modyfikacja kierunków semantycznych, agregowanie reprezentacji czy budowanie centroidów, takie operacje są użyteczne, a VQL mógłby zapewnić im czytelny zapis. Szczególnie ważna jest agregacja, na przykład `AVG(user_vector)` z `GROUP BY`. W klasycznym SQL średnia liczb jest zwykłą operacją statystyczną; w przestrzeni wektorowej średnia grupy może oznaczać centroid, czyli geometryczny środek klastra znaczeń. Otwiera to drogę do analiz na poziomie grup: można badać reprezentatywny wektor kategorii dokumentów, średni profil preferencji użytkowników, centroid tematów w zbiorze publikacji czy semantyczne przesunięcia między okresami, grupami lub autorami. Centroid nie jest „esencją” grupy w sensie filozoficznym, lecz matematycznym punktem centralnym, który dobrze użyty staje się cennym narzędziem eksploracji. Możliwość agregacji dowodzi, że bazy wektorowe mogą służyć nie tylko wyszukiwaniu, ale i analizie. Pozwala to pytać, które grupy dokumentów są do siebie zbliżone, które klastry są rozproszone, gdzie pojawiają się anomalie, jak zmienia się język organizacji w czasie lub jak różne wspólnoty mówią o tym samym problemie.

W takim ujęciu VQL nie byłby jedynie językiem dla chatbotów, lecz narzędziem analityki kulturowej, naukowej, ekonomicznej i instytucjonalnej. Dane semantyczne przestałyby być wyłącznie paliwem generacji, a stałyby się przedmiotem badania. Tu dochodzimy do społecznego znaczenia standaryzacji. Jeśli każdy dostawca baz wektorowych zachowa własny język, powstanie rynek pełen wysp. Na każdej wyspie można budować, ale trudno migrować; można optymalizować, lecz trudno porównywać; można eksperymentować, ale trudno uczyć wspólnej metodologii. VQL miałby działać jak most. Programista mógłby pisać logikę bardziej niezależną od konkretnego backendu, administrator tworzyć narzędzia cross-vendor, a badacz łatwiej opisywać eksperymenty. Organizacja mogłaby w ten sposób zmniejszyć ryzyko zamknięcia w jednym ekosystemie.

Standard nie jest zatem nudą komitetową; dobrze zaprojektowany standard to infrastruktura wolności. Oczywiście standaryzacja niesie ryzyka. Jeśli powstanie zbyt wcześnie, może utrwalić błędne abstrakcje. Jeśli zostanie przejęta przez kilku wielkich dostawców, może udawać neutralność, wspierając w praktyce ich interesy. Jeśli będzie zbyt ogólna, stanie się pustą fasadą, a jeśli zbyt szczegółowa – zdusi innowację. Historia technologii zna wiele standardów, które miały zbawić świat, a skończyły jako dokumenty PDF czytane tylko przez ludzi o wyjątkowo silnej psychice. Dlatego propozycja Borwankara, aby VQL wyrastał z praktyki open-source i dopiero z czasem dążył ku formalizacji, jest rozsądna.

Standaryzacja zapytań jako fundament edukacji i audytu

Najpierw żywy język, potem akademicka gramatyka. Odwrotna kolejność grozi technologiczną łaciną bez wiernych. VQL ma również znaczenie dla edukacji. Dziś osoba ucząca się baz wektorowych musi często jednocześnie zgłębiać embeddings, metryki, indeksy, API konkretnego narzędzia, ograniczenia filtrowania, strategie chunkingu i integrację z LLM. Wspólny język zapytań ułatwiłby naukę poprzez oddzielenie pojęć ogólnych od szczegółów implementacji. Student, programista, analityk czy badacz mógłby najpierw zrozumieć, czym jest similarity search, range search, hybrid search, threshold, metric, vector aggregation, a dopiero potem poznawać różnice między konkretnymi silnikami. To tak jak z SQL-em: można pojąć ideę SELECT i JOIN, zanim wejdzie się w dialekty PostgreSQL, MySQL czy SQLite.

Język kształtuje myślenie. Jeśli brakuje nam precyzyjnych słów i struktur składniowych dla operacji wektorowych, będziemy traktować je jak czarne skrzynki. Posiadając język, zaczynamy dostrzegać różnice: czy zapytanie wykorzystuje cosine similarity, czy L2? Czy wynik jest top-k, czy threshold-based? Czy filtr metadanych zastosowano przed wyszukiwaniem, czy po overfetch? Czy użyto wyszukiwania hybrydowego i jakie wagi przypisano komponentowi tekstowemu oraz wektorowemu? Czy operacja jest wsadowa? Czy agregujemy wektory?

Język czyni architekturę widzialną, a widzialność jest pierwszym warunkiem odpowiedzialności. Jest to szczególnie istotne w RAG-u, gdzie retrieval decyduje o tym, co model językowy uzna za kontekst. Użytkownik widzi końcową odpowiedź, lecz rzadko ma wgląd w zapytanie do bazy. Gdyby VQL lub podobny standard pozwolił zapisywać retrieval w sposób jawny, audytowalny i przenośny, znacząco poprawiłoby to kontrolę nad systemami generatywnymi. Można by porównywać strategie retrieval, testować progi, dokumentować metryki, wykrywać błędy i precyzyjniej wyjaśniać, dlaczego model odpowiedział w określony sposób. W epoce, w której AI wspiera decyzje prawne,

medyczne, finansowe czy administracyjne, niewidzialny retrieval stanowi zbyt poważne ryzyko. To jak wyrok bez akt sprawy: może brzmieć dostojnie, ale rozsądny człowiek zaczyna nerwowo szukać uzasadnienia.

VQL stawia również pytanie o relację między bazami relacyjnymi a wektorowymi. Nie należy przedstawiać ich jako wrogów. Relacyjność pozostaje niezastąpiona tam, gdzie liczy się spójność, transakcje i twarda struktura. Wektorowość dominuje natomiast w obszarze semantycznego podobieństwa, pracy z danymi nieustrukturyzowanymi i retrieval dla AI. Najbardziej dojrzałe systemy będą łączyć oba światy.

PostgreSQL z pgvector jest dobrym przykładem tej syntezy: klasyczna baza relacyjna otrzymuje typ wektorowy i indeksy podobieństwa. VQL jako idea może być językowym odpowiednikiem tego procesu: SQL-owa dyscyplina spotyka geometrię znaczenia. VQL nie jest więc tylko propozycją techniczną, lecz próbą ucywilizowania nowej warstwy cyfrowej rzeczywistości. Wraz z pojawieniem się nowego typu danych i infrastruktury decyzji, naturalnie rodzi się potrzeba norm, języka, audytu i przenośności. Bez tego technologia pozostaje plemienna: każdy dostawca ma własnego kapłana API, własne zaklęcia SDK i własny rytuał migracji. Z VQL mogłaby stać się bardziej republikańska: różne systemy, ale wspólne zasady rozmowy. Brzmi górnołotnie, ale czasem właśnie składnia jest początkiem cywilizacji.

Najważniejsze jest jednak to, że VQL ujawnia dojrzałość całego pola. Na początku rewolucji AI fascynowaliśmy się modelami, potem odkryliśmy embeddings, bazy wektorowe, FAISS, HNSW, RAG czy hybrydowe wyszukiwanie. Teraz pojawia się pytanie o język. To naturalna kolejność: najpierw narzędzia, potem praktyki, na końcu potrzeba gramatyki. Gdy praktyka dojrzewa, entuzjazm przestaje wystarczać; trzeba zacząć nazywać operacje i budować wspólną warstwę komunikacji. VQL jest zatem projektem przyszłościowym, lecz nie fantazyjnym.

Wynika on z realnych potrzeb: programistów pragnących pisać aplikacje niezależne od dostawcy; inżynierów danych, DevOps i SRE dążących do automatyzacji międzyplatformowej; oraz badaczy AI, którzy chcą eksperymentować na wektorach bez konieczności ciągłego pisania ciężkiego kodu pomocniczego.

VQL jako audytowalny pomost między intencją a danymi

Te trzy grupy obrazują skalę problemu. Nie chodzi tu jedynie o elegancję języka, lecz o koszty pracy, przenośność danych, automatyzację, proces nauki i realną kontrolę nad infrastrukturą. Aby

koncepcja VQL mogła się rozwinąć, musi zachować równowagę między prostotą a ekspresywnością. Zbyt uproszczony język nie obejmie realnych operacji wektorowych; zbyt złożony stanie się barierą w implementacji i nauce. Standard ten musiałby obsługiwać podstawowe operacje: similarity search, range search, hybrid search, batch search, a także metryki, progi, filtry metadanych, funkcje wektorowe i agregacje. Z czasem mógłby objąć reranking, wersjonowanie embeddings, deklarowanie modeli osadzania, strategie indeksów, polityki bezpieczeństwa oraz opis źródeł dla RAG. Jednak zbyt gwałtowna ekspansja groziłaby stworzeniem potwora składniowego. Język, podobnie jak instytucja, musi rosnąć, ale nie puchnąć.

Najciekawsze pytanie dotyczy przyszłej relacji VQL z systemami generatywnymi. Być może w dojrzałych architekturach RAG użytkownik nie będzie pisał kodu VQL bezpośrednio. Będzie zadawał pytania językiem naturalnym, a system przetłumaczy je na jawne i audytowalne zapytanie VQL. Wówczas model językowy przestanie być tylko generatorem odpowiedzi, a stanie się tłumaczem intencji użytkownika na formalny retrieval. Wymagałoby to jednak ścisłych mechanizmów kontroli: użytkownik lub administrator powinni mieć wgląd w to, jakie zapytanie wykonano, jakie filtry zastosowano, które metryki użyto i z jakich źródeł pobrano dane.

Język naturalny jest wygodny, ale formalny język jest sprawdzalny. Dojrzały system powinien łączyć oba te światy. VQL może zatem stać się warstwą pośrednią między człowiekiem, bazą wektorową a modelem językowym: człowiek formułuje problem, system przekształca go w zapytanie, baza zwraca semantycznie trafne fragmenty, a LLM syntetyzuje odpowiedź. Audyt pozwala sprawdzić, co wydarzyło się po drodze. To wizja AI jako infrastruktury odpowiedzialnej, a nie tylko efektywnej. Przyszłość bowiem nie należy do modeli, które najładniej mówią, lecz do systemów potrafiących wykazać, dlaczego mówią to, co mówią.

VQL domyka łuk całego artykułu. Zaczęliśmy od wektora jako listy liczb reprezentującej znaczenie. Przeszliśmy przez embeddings jako mechanizm osadzania sensu, FAISS jako silnik odnajdywania sąsiedztwa, SQLite-vss jako osobistą pamięć, PostgreSQL i pgvector jako naukową bazę semantyczną, RAG jako architekturę uziemienia LLM oraz lokalne modele jako projekt prywatności i kontroli. VQL jest próbą nadania temu wszystkiemu języka. A język nie jest dodatkiem – jest narzędziem porządku. Bez języka mamy sztuczki. Z językiem możemy budować instytucje.

Pozostaje pytanie końcowe: czy VQL stanie się przyszłym standardem, czy pozostanie jedynie ciekawą propozycją z książki? Tego dziś nie da się uczciwie przesądzić. Można jednak stwierdzić coś ostrożniejszego i zarazem mocniejszego: potrzeba, którą VQL nazywa, jest realna. Rynek baz wektorowych będzie potrzebował większej przenośności, jawności, audytowalności i wspólnej składni operacji. Niezależnie od tego, czy nazwiemy to VQL, rozszerzeniem SQL, standardem open-source czy warstwą pośrednią dla RAG – problem nie zniknie.

Bazy wektorowe jako mapy znaczeń a nie magazyny danych

Gdy dane zyskują geometrię znaczenia, konieczne staje się opracowanie gramatyki rozmowy z tą geometrią. Bazy wektorowe jako infrastruktura nowej epistemologii. Od wyszukiwarki do instytucji pamięci.

Po przejściu przez VQL krajobraz staje się jasny: wektory nie są techniczną ciekawostką, embeddings to coś więcej niż etap preprocessingu, FAISS nie jest wyłącznie biblioteką przyspieszającą obliczenia, RAG nie jest modnym dodatkiem do chatbota, a lokalne LLM nie są tylko zabawką dla entuzjastów open-source. Razem tworzą nową architekturę pracy z wiedzą. Przesuwa ona punkt ciężkości z przechowywania informacji na odnajdywanie sensu, z literalnego dopasowania na podobieństwo semantyczne, a model przestaje być samotnym generatorem, stając się uczestnikiem kontrolowanego rurociągu źródłowego. W klasycznej informatyce pytanie o dane brzmiało często: gdzie coś jest zapisane? W epoce baz wektorowych pytanie brzmi inaczej: co jest znaczeniowo blisko tego, czego szukam? Ta zmiana jest bardziej radykalna, niż może się wydawać.

Otwiera ona możliwość budowania systemów, które nie tylko katalogują dokumenty, lecz pomagają odkrywać relacje między nimi. Dokument przestaje być zamkniętym obiektem w folderze; staje się punktem w przestrzeni znaczeń, posiadającym sąsiadów, kierunki, odległości, klastry i potencjalne analogie. Baza wektorowa jest zatem mniej magazynem, a bardziej mapą. A dobra mapa nie tylko wskazuje położenie – pozwala zobaczyć układ terenu.

Jest to jednak mapa przybliżona. Wektor nie jest znaczeniem samym w sobie, lecz reprezentacją wytworzoną przez model na podstawie danych treningowych i procedury osadzania. Oznacza to, że bazy wektorowe dziedziczą zarówno siłę, jak i słabości modeli tworzących embeddings. Potrafią uchwycić podobieństwa niewidzialne dla wyszukiwania słów kluczowych, ale mogą też powielać stronniczości, gubić niuanse domenowe, spłaszczać różnice i błędnie łączyć teksty, które są podobne językowo, lecz odmienne merytorycznie. Notatka źródłowa podkreśla te ograniczenia: spadek jakości w niszowych domenach, koszty obliczeniowe większej wymiarowości oraz dziedziczenie uprzedzeń z danych treningowych.

Dlatego pierwszą zasadą dojrzałego użycia baz wektorowych powinien być realizm. Nie wolno twierdzić, że system „rozumie” dokumenty w sensie ludzkim. Trafniej jest powiedzieć, że tworzy matematyczne reprezentacje, które często zachowują użyteczne relacje semantyczne. To mniej efektowne, ale bardziej prawdziwe. W tej dziedzinie precyzja języka jest zabezpieczeniem przed bałwochwalstwem technologicznym. Gdy mówimy, że baza „rozumie sens”, łatwo zapominamy, że

sens ten został przefiltrowany przez model, metrykę, chunking, próg podobieństwa i indeks. W praktyce „rozumienie” jest tu całym łańcuchem operacji, a nie pojedynczym aktem inteligencji.

Druga zasada dotyczy źródłowości. System RAG jest tyle wart, ile dokumenty, które potrafi odnaleźć i poprawnie wykorzystać. Jeśli korpus jest ubogi, przestarzały, niespójny lub źle oczyszczony, model wygeneruje odpowiedź na kruchych podstawach. Jeśli retrieval pobierze fragmenty nieadekwatne, nawet najlepszy LLM będzie składał wnioski z materiału przypadkowego. Jeżeli prompt pozwoli modelowi dopowiadać fakty spoza źródeł, system znów zacznie zachowywać się jak klasyczny generator halucynacji, tyle że z ozdobnym szyldem „RAG”. Notatka źródłowa wielokrotnie akcentuje konieczność wymuszania cytowań, pracy wyłącznie na dostarczonym kontekście i przyznawania niewiedzy, gdy kontekst nie zawiera odpowiedzi.

Trzecia zasada dotyczy hybrydowości. Najbardziej obiecujące systemy nie będą ani czysto relacyjne, ani czysto wektorowe. Będą łączyć twardą strukturę metadanych z miękkim podobieństwem semantycznym. Będą potrafiły wskazać: szukaj sensu, ale tylko w dokumentach z tej daty, tego autora, tej kategorii i typu, z tym poziomem uprawnień i powyżej określonego progu trafności. To połączenie jest kluczowe dla profesjonalnych zastosowań. Sama semantyka może być zbyt rozlana, sama struktura – zbyt ślepa. Dopiero razem tworzą narzędzie, które potrafi pracować z rzeczywistą wiedzą, a nie tylko z jej formalnym cieniem.

PostgreSQL z pgvector jest w tym kontekście czymś więcej niż praktycznym wyborem technologicznym; jest symbolem syntezy. Relacyjna tradycja baz danych wnosi rygor, transakcyjność, metadane, klucze, filtry i dojrzałość administracyjną. Warstwa wektorowa dodaje zdolność pracy na nieustrukturyzowanym sensie. Takie połączenie może być szczególnie ważne dla instytucji: uczelni, kancelarii, fundacji, redakcji, administracji, laboratoriów i przedsiębiorstw. Organizacje nie potrzebują wyłącznie chatbotów.

Bazy wektorowe jako lekarstwo na amnezję instytucjonalną

Organizacje potrzebują pamięci, która potrafi odszukać wcześniejsze argumenty, decyzje, analizy, analogiczne przypadki i powiązane dokumenty. Potrzebują systemów, które nie tylko generują wypowiedź, ale wydobywają własne doświadczenie z archiwum. Prowadzi to do praktycznej obserwacji: w wielu instytucjach największym problemem nie jest brak wiedzy, lecz brak dostępu do tej już wytworzonej. Dokumenty, analizy i uzasadnienia istnieją; fragmenty strategii są zapisane. Ktoś przygotował stanowisko, zebrał dane czy opisał precedens, ale plik ugrzązł w folderze, mailu,

komunikatorze, dysku wspólnym lub w prywatnej pamięci pracownika. Baza wektorowa może stać się narzędziem walki z taką amnezją. Nie wytworzy ona sama instytucjonalnej mądrości, ale zapobiegnie sytuacji, w której jest ona regularnie gubiona pod warstwą nazw plików i zmian personalnych. W tym miejscu technologia spotyka socjologię instytucji.

Organizacja, która nie potrafi odnaleźć własnej wiedzy, zachowuje się jak człowiek z uszkodzoną pamięcią roboczą: reaguje impulsywnie, powtarza błędy, odtwarza te same spory i uzależnia się od pojedynczych osób. System semantycznej pamięci może zmienić tę dynamikę. Pozwala nowemu pracownikowi prześledzić historię argumentu, zespołowi odnaleźć podobną sprawę sprzed dwóch lat, redakcji przywrócić wcześniejszą linię interpretacyjną, a fundacji zachować ciągłość pracy legislacyjnej.

To nie jest futurystyczna ekstrawagancja, lecz higiena poznawcza organizacji. Wymaga ona jednak dyscypliny. Nie wystarczy wrzucić dokumentów do bazy i polecić AI działanie. Konieczne jest określenie, które pliki są aktualne, a które archiwalne; które robocze, a które zatwierdzone; które poufne, a które mogą być cytowane jako materiał pomocniczy. Należy zachować metadane: daty, autorów, wersje, relacje między dokumentami oraz kontekst ich powstania. W przeciwnym razie baza wektorowa może mieszać wersję roboczą z finalną, nieaktualny argument z obowiązującym stanowiskiem lub hipotezę z konkluzją. AI nie posiada naturalnego szacunku dla obiegu dokumentów – trzeba jej go narzucić strukturą.

Jest to szczególnie istotne w dziedzinach o wysokiej odpowiedzialności: prawie, medycynie, finansach, administracji czy nauce. Tam odpowiedź modelu musi być nie tylko poprawnie brzmiąca, ale przede wszystkim audytowalna. Powinna wskazywać źródła, rozróżniać poziom pewności i identyfikować sprzeczności, nie przekraczając materiału dowodowego. RAG może to wspierać, lecz nie gwarantuje tego automatycznie; architektura musi być zaprojektowana z myślą o odpowiedzialności. W przeciwnym razie powstanie system, który z wielką płynnością będzie produkował półprawdy. A półprawda w języku eksperckim jest jak trucizna w kryształowej karafce: elegancka, ale nadal trująca.

Kluczową kwestią pozostaje lokalność i prywatność. Przetwarzanie wrażliwych danych na własnym sprzęcie, brak wysyłania dokumentów do sieci oraz redukcja kosztów operacyjnych sprawiają, że lokalne LLM stają się niezbędne. Nie jest to prosta negacja chmury, która bywa efektywna i wygodna, lecz konieczna alternatywa tam, gdzie kontrola i autonomia są ważniejsze niż maksymalna elokwencja największego modelu. Przyszłość będzie zapewne hybrydowa: podział zadań między infrastrukturę lokalną, prywatną i modele publiczne.

Dojrzałość polega na rozróżnianiu, nie na sloganie. Istotny jest również koszt poznawczy. Bazy wektorowe ułatwiają wyszukiwanie, ale mogą wytwarzać iluzję łatwego rozumienia. Użytkownik otrzymuje odpowiedź szybciej i bardziej syntetycznie niż przy ręcznym przeszukiwaniu dokumentów, co może osłabić jego czujność. Jeśli system brzmi pewnie, a wyniki są zawsze płynne, użytkownik może przestać sprawdzać źródła lub nie zauważyć słabego retrievalu. Wtedy AI przestaje być narzędziem wzmacniania rozumu, a staje się elegancką maszyną zastępczego rozumienia.

Ryzyko to rośnie w środowisku zawodowym, gdzie presja czasu sprzyja delegowaniu myślenia maszynie. RAG może skrócić drogę do źródła, ale nie znosi obowiązku krytyki. Model pomaga odnaleźć fragmenty, lecz człowiek musi ocenić ich wagę; model buduje syntezę, lecz ekspert sprawdza brak kontrargumentów; model streszcza orzeczenie, lecz prawnik musi zrozumieć ratio decidendi. W przeciwnym razie zamiast społeczeństwa wiedzy stworzymy społeczeństwo streszczeń. A streszczenie bez odpowiedzialności to fast food dla intelektu: szybkie, miękkie i po pewnym czasie szkodliwe.

Wreszcie, istotna staje się multimodalność. Modele potrafiące tworzyć kompatybilne reprezentacje dla tekstu, dźwięku, kodu, obrazów czy tabel otwierają przed organizacjami ogromne możliwości zarządzania zróżnicowanymi typami danych.

Multimodalność i standardy jako warunek audytowalności

Publikacja naukowa to nie tylko akapity, ale również wykresy, schematy, wzory i tabele. Dokument organizacyjny może zawierać skany, fotografie, diagramy procesów czy prezentacje, a nagranie – wiedzę, która nigdy nie została spisana. Jeśli wszystkie te formy uda się osadzić w powiązanych przestrzeniach semantycznych, systemy wiedzy staną się znacznie bogatsze. Jednocześnie jednak multimodalność potęguje problem weryfikacji.

Gdy model cytuje zdanie z tekstu, użytkownik może je po prostu przeczytać. Gdy odwołuje się do wykresu, trzeba mieć pewność, że prawidłowo odczytał osie, skalę, legendę i relacje. Przy interpretacji obrazu należy sprawdzić, czy nie pomylił podobieństwa wizualnego z znaczeniowym, a w przypadku tabeli – czy zachował strukturę danych. Multimodalny RAG będzie zatem wymagał nowych form cytowania: nie tylko numeru dokumentu, ale konkretnego regionu obrazu, komórki tabeli, fragmentu osi wykresu czy znacznika czasu nagrania. W przeciwnym razie otrzymamy piękne interpretacje nie wiadomo czego. A "nie wiadomo czego" to kiepski fundament wiedzy.

Siódma zasada dotyczy standardów. Koncepcja VQL pokazuje, że pole baz wektorowych dojrzewa do wspólnej gramatyki. Celem VQL jest ustandaryzowanie rozdrobnionego rynku poprzez stworzenie interfejsu niezależnego od dostawcy. To nie jest techniczny luksus. Bez standardów organizacje zostaną zamknięte w ekosystemach konkretnych producentów, programiści będą zmuszeni przepisywać logikę retrieval dla różnych baz, badacze napotkają trudności z powtarzalnością eksperymentów, a użytkownicy nie będą rozumieli, jak naprawdę działa ich system. Standardy bywają nudne tylko dla tych, którzy nigdy nie płacili rachunku za migrację.

VQL ma szczególne znaczenie, ponieważ dotyczy warstwy niewidzialnej dla większości osób. Użytkownik widzi odpowiedź modelu, lecz rzadko wie, jakie zapytanie zostało wykonane, jaką zastosowano metrykę i próg podobieństwa, czy użyto wyszukiwania hybrydowego, ile wyników pobrano, czy poddano je rerankingowi oraz czy filtr metadanych zadziałał przed, czy po wyszukiwaniu. To właśnie te decyzje kształtują kontekst, który otrzymuje model. Bez jawnej gramatyki proces retrieval pozostaje czarną skrzynką, którą VQL mógłby pomóc otworzyć.

Ósma zasada dotyczy odpowiedzialności za ranking. Wyszukiwanie semantyczne zwraca wyniki uporządkowane według podobieństwa, jednak ranking nie jest neutralny. To, co znajduje się wysoko, ma większą szansę zostać wykorzystane przez model i zauważone przez człowieka; to, co spadnie niżej, może zostać całkowicie pominięte.

Systemy wiedzy jako narzędzia sprawczości, a nie protezy myślenia

W klasycznych wyszukiwarkach problem rankingu od dawna ma wymiar polityczny, ekonomiczny i kulturowy. W bazach wektorowych powraca on w nowej formie: ranking podobieństwa może bezpośrednio kształtować odpowiedź AI. Jeśli system regularnie pobiera określony typ dokumentów jako „najbardziej podobne”, ryzykuje utrwaleniem pewnego obrazu wiedzy przy jednoczesnej marginalizacji innych głosów. Dlatego metryki, progi i strategie retrieval powinny być testowane nie tylko pod kątem technicznym, ale i epistemicznym.

Dziewiąta zasada dotyczy aktualizacji. Wektorowa baza wiedzy nie jest raz na zawsze gotowa; dokumenty ewoluują, a wiedza się starzeje. Modele embeddings są aktualizowane, a ich nowe wersje mogą generować inne reprezentacje, co wymusza przebudowę indeksów i archiwizację starych treści. Jeśli system nie zarządza cyklem życia danych, zaczyna udawać aktualność. Jest to szczególnie groźne w dziedzinach krytycznych czasowo: w prawie, regulacjach, medycynie, technologiach czy polityce publicznej. RAG oparty na nieaktualnym korpusie jest jak doradca, który

czyta gazety sprzed pięciu lat i mówi z pewnością prezesa na konferencji. Efekt bywa przekonujący, lecz kalendarz ma inne zdanie.

Dziesiąta zasada dotyczy człowieka. Wszystkie opisane technologie — embeddings, FAISS, pgvector, sqlite-vss, RAG, lokalne LLM czy VQL — powinny zwiększać sprawczość użytkownika, a nie ją odbierać. Dobry system wiedzy pozwala szybciej dotrzeć do źródeł, sprawniej porównywać dokumenty, pamiętać wcześniejsze argumenty, wykrywać analogie i tworzyć syntezy, rozpoznając jednocześnie luki w informacji. Zły system zastępuje myślenie wygodną odpowiedzią. Różnica między nimi nie tkwi wyłącznie w technologii, lecz w kulturze użycia, procedurach, audycie, edukacji oraz gotowości użytkownika do zadawania niewygodnych pytań. Technologia może wzmocnić rozum, ale może go również uspić; jak każde profesjonalne narzędzie, wymaga charakteru.

Warto w tym miejscu wrócić do pojęcia „drugiego mózgu”, obecnego w opisach osobistych systemów zarządzania wiedzą. Metafora ta jest atrakcyjna, lecz wymaga ostrożności. Drugi mózg nie może stać się protezą lenistwa; powinien być rozszerzeniem pamięci roboczej, instrumentem odnajdywania i porządkowania informacji. Ludzki mózg nie jest bowiem tylko magazynem danych, lecz ośrodkiem oceny, intencji, wartości i doświadczenia. System wektorowy wspiera pamięć, ale nie zastąpi sądu. Jeśli o tym zapomnimy, drugi mózg szybko stanie się pierwszym alibi.

Ta uwaga prowadzi do wymiaru aksjologicznego. Bazy wektorowe i RAG nie są neutralne społecznie. Mogą służyć prywatności lub nadzorowi; wzmacniać małe organizacje albo pogłębiać zależność od wielkich platform. Mogą demokratyzować dostęp do wiedzy lub tworzyć nowe formy zamknięcia, wspierać rzetelność źródłową bądź generować jeszcze bardziej przekonujące halucynacje. Mogą pomagać w nauce, a mogą przyspieszać powierzchowne streszczanie treści bez ich faktycznego czytania.

Każda infrastruktura wiedzy jest również infrastrukturą władzy. Kto kontroluje pamięć, retrieval i ranking, ten wpływa na to, co zostanie uznane za relewantne. Dlatego lokalne systemy open-source mają znaczenie wykraczające poza technikę. Pozwalają eksperymentować, kontrolować dane i budować rozwiązania dopasowane do mniejszych wspólnot wiedzy. Nie wszyscy muszą korzystać z tej samej chmurowej pamięci, nie każda organizacja musi oddawać swoje archiwum największemu dostawcy, a nie każdy badacz musi ufać zamkniętemu pipeline'owi. Oczywiście open-source nie jest wolny od błędów czy stronniczości, ale daje możliwość wglądu i niezależnej weryfikacji. W epoce AI są to wartości nie techniczne, lecz obywatelskie. Jednocześnie należy unikać romantyzowania samodzielności — własny system wymaga kompetencji.

Dojrzała architektura AI jako system warstwowy

Źle skonfigurowany lokalny RAG może być gorszy niż dobrze zabezpieczona usługa chmurowa. Lokalna baza bez kopii zapasowych to katastrofa czekająca na awarię dysku, a własny model bez aktualizacji z czasem zacznie odstawać. Własny pipeline pozbawiony dokumentacji stanie się prywatnym labiryntem twórcy, natomiast autonomia bez procedury staje się chaosem z dobrymi intencjami. Dlatego prawdziwa dojrzałość polega na łączeniu niezależności z profesjonalizacją. Książka Borwankarze ma zatem wartość nie tylko techniczną, lecz także pedagogiczną, gdyż uczy myślenia systemowego.

Pokazuje ona, że nie istnieje jeden cudowny element. Wektor wymaga embeddings; te z kolei wymagają dobrego modelu i preprocessingu. Wyszukiwanie opiera się na indeksie i metryce, a RAG wymaga retriewalu, promptu, cytowania oraz kontroli temperatury. Lokalność wymusza dbałość o sprzęt i procedury, natomiast VQL wymaga wspólnoty praktyki. Każdy z tych elementów jest częścią większej całości, co stanowi cenną lekcję dla epoki, która lubi magiczne skróty. AI nie jest bowiem skrótem od myślenia, lecz powiększeniem skutków naszego myślenia – zarówno tych dobrych, jak i złych.

Jeżeli spojrzymy na całość z perspektywy historii idei, bazy wektorowe przypominają kolejny etap długiej walki człowieka z nadmiarem informacji. Biblioteka rozwiązywała problem przechowywania tekstów, katalog – problem ich odnajdywania, a indeks – kwestię lokalizacji pojęć. Wyszukiwarka internetowa zmierzyła się ze skalą sieci, natomiast baza wektorowa próbuje rozwiązać problem sensu: jak znaleźć coś, co jest podobne znaczeniowo, choć niekoniecznie słownie. RAG dodaje do tego problem odpowiedzi: jak z odnalezionych fragmentów zbudować wypowiedź wiernie związaną ze źródłami. VQL wprowadza natomiast problem języka: jak opisywać takie operacje w sposób wspólny i audytowalny.

To nie oznacza, że dawne narzędzia stają się przestarzałe. Przeciwnie: najlepsze systemy będą warstwowe. Katalog, metadane, pełnotekstowe wyszukiwanie oraz SQL nadal pozostają potrzebne. Wektory dodają nową warstwę, ale nie unieważniają poprzednich. Technologiczny postęp często przypomina wojnę pokoleń narzędzi, jednak dojrzała infrastruktura kształtuje się raczej na wzór dobrze zorganizowanego archiwum, gdzie stare i nowe metody współpracują, gdyż każda odpowiada na inny rodzaj pytania.

W praktyce profesjonalny system wiedzy powinien umożliwiać kilka trybów pracy. Użytkownik powinien móc wyszukiwać: - literalnie, gdy potrzebuje konkretnej frazy, - semantycznie, gdy szuka podobnego sensu, - hybrydowo, gdy potrzebuje obu tych funkcji naraz, - strukturalnie, gdy filtruje po metadanych, - źródłowo, gdy wymaga cytowań, - porównawczo, gdy zestawia dokumenty, -

historycznie, gdy bada zmiany w czasie, - analitycznie, gdy grupuje i agreguje embeddings, - generatywnie, gdy chce otrzymać syntezę. Każdy z tych trybów niesie własne ryzyka i posiada specyficzne zastosowania.

Systemy AI muszą być sprawdzalne, a nie magiczne

Użytkownik powinien je rozumieć, a system – ujawniać. Największym błędem byłoby potraktowanie baz wektorowych jako kolejnego etapu automatyzacji bez zmiany kultury pracy. Jeśli organizacja operuje na chaotycznych dokumentach, przy braku wersjonowania, niejasnych uprawnieniach i słabej kulturze źródłowej, RAG nie stanie się cudownym lekarstwem; może wręcz przyspieszyć chaos.

Najpierw należy uporządkować korpus: ustalić status dokumentów, zadbać o metadane, opracować politykę prywatności i określić standard odpowiedzi. Dopiero wtedy AI może wzmocnić instytucję. W przeciwnym razie będzie to jak zatrudnienie genialnego archiwisty do magazynu bez półek, etykiet i światła. Będzie pracował, ale głównie nad własnym rozpaczą. Warto więc myśleć o bazach wektorowych nie jako o produkcie, lecz jako o praktyce. Obejmuje ona przygotowanie danych, wybór modelu embeddings, dobór metryki, indeksowanie, testy jakości, progi podobieństwa, strategie hybrydowe, decyzję między lokalnością a chmurą, prompt engineering, cytowanie, audyt, aktualizacje oraz szkolenie użytkowników. Dopiero suma tych elementów tworzy system godny zaufania. Sam zakup narzędzia niczego nie gwarantuje – w świecie AI wdrożenie bez metodologii przypomina zakup mikroskopu przez kogoś, kto nie wie, czym jest próbka. Coś zobaczy, ale nie wiadomo co.

Na tym tle VQL jawi się jako element profesjonalizacji. Jeśli zapytania wektorowe staną się czytelne, standaryzowane i przenośne, łatwiej będzie budować procedury jakościowe, porównywać wyniki między systemami oraz uczyć nowych użytkowników. Ułatwi to audyt retrieval w RAG-u, migrację danych oraz odróżnienie błędu modelu od błędu wyszukiwania. Standaryzacja nie rozwiąże wszystkich problemów, ale może zmniejszyć dzisiejszą mgłę implementacyjną. Ta mgła jest bowiem bardzo wygodna dla dostawców i bardzo kosztowna dla użytkowników.

W perspektywie długoterminowej możemy wyobrazić sobie systemy, w których użytkownik zadaje pytanie w języku naturalnym, model tłumaczy je na formalne zapytanie wektorowe, baza zwraca wyniki, reranker porządkuje dowody, a LLM generuje odpowiedź. Użytkownik widziałby wówczas nie tylko wynik, lecz całą ścieżkę: zapytanie, metrykę, źródła, progi, fragmenty i poziom pewności. Byłby to istotny krok od AI jako czarnej skrzynki ku narzędziu podlegającemu kontroli. Nie chodzi

o to, by każdy użytkownik codziennie analizował techniczne szczegóły, lecz by system mógł je ujawnić w razie potrzeby. Jawność nie musi być natrętna, ale musi istnieć.

W takim modelu AI staje się mniej „magiczna”, a bardziej godna zaufania. Magia jest atrakcyjna w reklamie, jednak w instytucjach potrzebujemy procedur. Największe systemy przyszłości nie będą tylko najbardziej kreatywne – będą przede wszystkim sprawdzalne. Potrafią nie tylko odpowiedzieć, lecz wskazać drogę od pytania do źródła; odróżnić brak danych od braku odwagi i jasno zakomunikować: tego nie wiem, tego nie ma w korpusie, to jest sprzeczne lub wymaga aktualizacji.

Taka AI może być mniej efektowna na pokazie, ale znacznie bardziej użyteczna w pracy. Całą opowieść o wektorach można streścić w jednym zdaniu: przyszłość AI zależy nie tylko od modeli, które mówią, lecz od pamięci, z której mówią. Jeśli pamięć jest cudza, nieprzejrzysta lub źle zorganizowana, odpowiedź będzie podejrzana, nawet jeśli brzmi pięknie. Jeśli zaś jest własna, uporządkowana i audytowalna, model staje się realnym asystentem wiedzy. Bazy wektorowe są mniej widowiskowe niż chatboty, ale być może ważniejsze. Chatbot jest twarzą, baza wiedzy – kręgosłupem. A cywilizacja, która inwestuje tylko w twarz, kończy zwykle w gabinecie medycyny estetycznej, nie w laboratorium poznania.

Ta część prowadzi nas ku syntezie. Należy zebrać najważniejsze konsekwencje dla programistów, badaczy, instytucji i kultury wiedzy, by pokazać, dlaczego Borwankarowska wizja jest jednocześnie praktyczna i normatywna. Uczy ona budować systemy, ale sugeruje też, jakiej AI powinniśmy chcieć: nie sztucznej inteligencji jako centralnego megafonu platform, lecz systemów wiedzy lokalnych tam, gdzie trzeba, hybrydowych tam, gdzie warto, źródłowych tam, gdzie wymaga tego prawda i standaryzowanych tam, gdzie od standardu zależy wolność użytkownika.

Jaka sztuczna inteligencja zasługuje na zaufanie? Po przejściu przez wektory, embeddings, FAISS, SQLite-vss, PostgreSQL z pgvector, RAG, lokalne LLM i VQL można postawić pytanie najważniejsze: nie tylko jak te technologie działają, lecz jaką wizję AI wspierają. Technologia nigdy nie jest jedynie zestawem bibliotek i indeksów – jest projektem organizacji wiedzy, a taki projekt zawsze niesie konsekwencje społeczne, ekonomiczne i aksjologiczne. Kto kontroluje dane, kontroluje pamięć. Kto kontroluje pamięć, wpływa na odpowiedź. A kto wpływa na odpowiedź, współdecyduje o tym, co użytkownik uzna za prawdopodobne, relewantne i godne dalszego myślenia.

AI jako system odpowiedzialności a nie retoryki

Książka Borwankara w swojej zrekonstruowanej formie nie jest opowieścią o pojedynczej technologicznej sztuczce. To opowieść o całym stosie poznawczym: od wektorowej reprezentacji znaczenia, przez szybkie wyszukiwanie podobieństwa, aż po systemy RAG, które ograniczają halucynacje modeli językowych poprzez podłączenie ich do kontrolowanego korpusu źródeł. Główna teza brzmi zatem następująco: nowoczesna AI nie powinna być rozumiana jako sam model językowy, lecz jako system złożony z modelu, pamięci, wyszukiwarki, metadanych, promptu, mechanizmów cytowania, polityki prywatności i standardów zapytań. To przesunięcie jest zasadnicze. W pierwszej fali fascynacji generatywną AI użytkownicy skupiali się na samej odpowiedzi: czy jest płynna, efektowna, brzmi ludzko lub pisze elegancko? Dojrzałe pytanie brzmi jednak inaczej: skąd odpowiedź wie to, co twierdzi? Czy opiera się na źródłach i czy były one aktualne? Czy retrieval pobrał właściwe fragmenty? Czy model odróżnił fakt od interpretacji i potrafił powiedzieć „nie wiem”? Czy cytowania są powiązane z twierdzeniami, a dokumenty prywatne pozostały prywatne? Czy użytkownik może prześledzić drogę od pytania do odpowiedzi?

Właśnie dlatego RAG jest czymś więcej niż techniką; stanowi próbę przywrócenia odpowiedzialności wypowiedzi generowanej przez model. Model językowy bez retrieval przypomina mówcę o znakomitej dykcji i ogromnej pamięci wzorców, lecz z niepokojącą skłonnością do pewności siebie. RAG każe temu mówcy usiąść przy aktach sprawy. Nie czyni go nieomylnym, ale odbiera mu najgroźniejszą swobodę: możliwość produkowania odpowiedzi bez oparcia w dostarczonym materiale. Dobrze skonstruowany system RAG powinien wymuszać użycie wyłącznie podanego kontekstu, cytowanie źródeł i przyznanie się do niewiedzy, jeśli odpowiedź nie znajduje się w materiale. Zaufanie do AI nie powinno więc wynikać z tonu odpowiedzi, lecz z architektury odpowiedzialności.

To różnica między retoryką a epistemologią. Retoryka pyta, czy odpowiedź brzmi przekonująco; epistemologia pyta, czy ma ona podstawę. Retoryka może zachwycić, ale Epistemologia ma obowiązek przeszkadzać zachwytowi, gdy ten za bardzo się rozpędza. W świecie generatywnej AI sceptycyzm nie jest zrzędzeniem – jest pasem bezpieczeństwa. Lepiej zapiąć go przed wypadkiem niż po przeczytaniu pięknie napisanego nonsensu. Pierwszą konsekwencją dla programistów jest konieczność myślenia systemowego. Nie wystarczy „podłączyć model”. Trzeba zrozumieć cały łańcuch: pochodzenie dokumentów, ich czyszczenie i podział na chunki, model tworzący embeddings, metrykę podobieństwa, indeks przyspieszający wyszukiwanie oraz próg odcinający słabe wyniki. Należy wiedzieć, czy stosuje się wyszukiwanie hybrydowe, czy wyniki są rerankowane, jak budowany jest prompt, czy wymuszane są cytowania i jakie dane mogą opuścić infrastrukturę

użytkownika. Każdy z tych elementów może stać się źródłem błędu, ale każdy może również stać się miejscem poprawy jakości.

AI jako narzędzie wymagające nowej alfabetyzacji i standardów

Drugą konsekwencją dla badaczy jest zmiana sposobu pracy z literaturą. Wyszukiwarka semantyczna nie zastępuje krytycznego przeglądu źródeł, lecz może go przyspieszyć i pogłębić. Pozwala odnajdywać prace podobne znaczeniowo, nawet jeśli posługują się innym słownictwem; umożliwia przejście od ogólnych abstraktów do konkretnych fragmentów sekcji oraz wykrywanie nieoczywistych związków między dyscyplinami. Dzięki temu badacz nie utknie w jednym słowniku pojęć. Wymaga to jednak ostrożności: podobieństwo semantyczne nie jest dowodem prawdziwości, a ranking wyników nie zastępuje recenzji naukowej. Wektor wskazuje jedynie sąsiedztwo; ocena metody, argumentacji i siły dowodu wciąż spoczywa na barkach człowieka.

Trzecią konsekwencją, tym razem dla instytucji, jest konieczność uporządkowania własnej pamięci. Bazy wektorowe niosą ogromny potencjał tam, gdzie organizacje toną w dokumentach, nie potrafiąc wydobyć z nich wiedzy. Fundacje, kancelarie, uczelnie, redakcje, administracja, firmy technologiczne, związki zawodowe czy think tanki mogą dzięki nim odzyskiwać wcześniejsze stanowiska, analogiczne sprawy, zapomniane argumenty oraz archiwalne wersje raportów i materiałów źródłowych. Warunkiem jest jednak dojrzała kultura danych: rzetelne wersjonowanie, metadane, status dokumentów, uprawnienia, aktualność, polityka prywatności i procedury cytowania. AI nie uporządkuje instytucji, która sama nie wie, co jest dokumentem finalnym, a co notatką z pola bitwy.

Czwartą konsekwencją jest kwestia lokalności. Borwankar wyraźnie wskazuje na sens wdrażania lokalnych modeli LLM i baz danych w kontekście prywatności, kontroli oraz redukcji kosztów operacyjnych. Nie oznacza to, że chmura jest z natury zła, lecz że nie powinna być domyślnym odruchem przy każdym rodzaju wiedzy. Usługi zewnętrzne są uzasadnione przy pracy na danych publicznych lub mniej wrażliwych, wymagających największych modeli. Jeśli jednak operujemy na dokumentach prywatnych, badawczych, prawnych czy strategicznych, lokalny RAG staje się poważną alternatywą. Wygoda jest ważna, ale czasem wygoda jest tylko elegancko zapakowaną zależnością.

Piątą konsekwencją jest potrzeba standardów. VQL pojawia się w tej opowieści jako projekt przyszłościowy i eksperymentalny, jeszcze niegotowy na miano oficjalnego standardu, lecz trafnie

diagnozuje realny problem: rozdrobnienie rynku baz wektorowych i brak wspólnej gramatyki operacji podobieństwa. Bez takiej gramatyki programiści pozostaną uzależnieni od dialektów dostawców, badacze napotkają trudności z powtarzalnością wyników, a organizacje zapłacą wysoką cenę za każdą migrację. Standard nie jest nudą – jest formą wolności od kaprysu narzędzia, o ile oczywiście nie zostanie napisany przez największych graczy wyłącznie pod ich własne interesy.

Szóstą konsekwencją jest nowa rola użytkownika. Nie powinien on być biernym odbiorcą odpowiedzi, lecz krytycznym operatorem systemu wiedzy. Musi pytać o źródła, zakres korpusu, daty, progi podobieństwa, ograniczenia modelu, jakość retrieval oraz różnicę między faktem a interpretacją. To wymaga edukacji, gdyż AI nie zwalnia z kompetencji, lecz przesuwając je na nowy poziom.

Dawniej trzeba było umieć szukać w katalogu, indeksie, bazie prawniczej lub bibliografii. Dziś należy ocenić, czy system semantyczny poprawnie wytyczył drogę od pytania do odpowiedzi. Nadchodząca alfabetyzacja AI nie polega zatem na umiejętności wpisania efektownego promptu – to za mało. Prawdziwa alfabetyzacja obejmuje rozumienie embeddings, chunkingu, retrieval, cytowania, temperatury, ograniczeń modeli oraz prywatności danych.

AI jako system audytowalny i ekonomicznie niezależny

Nie każdy musi być inżynierem, tak jak nie każdy kierowca musi być mechanikiem. Jednak profesjonalny użytkownik powinien mieć świadomość, że samochód posiada hamulce, paliwo, silnik i martwe pole. W przypadku AI to właśnie martwe pole bywa szczególnie niebezpieczne, gdyż model potrafi skutecznie zamaskować je poprawnie sformułowanym akapitem.

Siódma konsekwencja dotyczy ekonomii wiedzy. Bazy wektorowe mogą obniżyć koszty odnajdywania informacji, ale jednocześnie tworzą nowe zależności rynkowe. Jeśli pamięć organizacji zostanie zamknięta w cudzym systemie, migracja stanie się trudna. Jeśli embeddingi zostaną wygenerowane przez model dostawcy bez jasnej dokumentacji, ich porównywalność może być ograniczona. Gdy retrieval pozostaje ukryty, audyt staje się kosztowny, a przy działaniu modeli wyłącznie przez API koszty rosną wraz ze skalą użycia. Architektura AI jest więc decyzją ekonomiczną, nie tylko techniczną. Kto tego nie widzi na początku, zwykle zobaczy na fakturze.

Ósma konsekwencja odnosi się do kultury źródłowej. RAG może przywrócić znaczenie dokumentu, lecz stanie się to możliwe tylko wtedy, gdy użytkownicy i organizacje będą wymagać cytowań, weryfikacji i jawnego rozróżniania podstawy odpowiedzi. Bez tych wymagań RAG pozostanie jedynie marketingowym skrótem. Otrzymamy odpowiedzi „oparte na wiedzy”, ale bez realnej

możliwości sprawdzenia, jaka konkretnie wiedza została wykorzystana i w jaki sposób. Byłaby to porażka nie technologii, lecz kultury. Narzędzie da się zbudować; trudniej zbudować obyczaj wymagania dowodu.

Dziewiąta konsekwencja dotyczy multimodalności. Przyszłe systemy wiedzy będą operować nie tylko na tekście, ale także na obrazach, tabelach, wykresach, dźwięku, wideo i kodzie. Modele multimodalne mogą tworzyć wspólne reprezentacje dla różnych typów danych, co pozwoli lepiej analizować elementy publikacji naukowych, które sprawiają trudność klasycznym parserom. Otwiera to ogromne możliwości, ale wymaga nowych standardów weryfikacji. Cytowanie akapitu to jedno; cytowanie fragmentu wykresu, tabeli lub obrazu to zupełnie inna kwestia. Im bogatsza modalność, tym większa potrzeba precyzyjnego wskazywania źródła twierdzenia.

Dziesiąta konsekwencja jest najbardziej ogólna: AI godna zaufania będzie mniej widowiskowa, a bardziej sprawdzalna. Nie musi zawsze odpowiadać natychmiast i efektownie; powinna natomiast umieć odpowiedzieć ostrożnie, wskazać źródła i przyznać brak danych. Powinna rozróżniać, co wynika bezpośrednio z dokumentu, a co jest jedynie interpretacją. Musi pozwalać na audyt oraz chronić prywatność tam, gdzie jest ona warunkiem zaufania. Wreszcie powinna być standaryzowana wszędzie tam, gdzie brak standardu czyni użytkownika zakładnikiem.

Wartość AI tkwi w audytowalności i pamięci, nie w elokwencji

Innymi słowy: dobra AI to nie ta, która najładniej mówi. Dobra AI jest tą, której drogę rozumowania można odtworzyć bez seansu spirytystycznego z czarną skrzynką. Spójrzmy zatem na sam początek: na wektor jako listę liczb. Na pierwszy rzut oka trudno wyobrazić sobie mniej humanistyczny obiekt. Lista floatów. Kilkaset wymiarów. Geometryczne sąsiedztwa, metryki, indeksy. A jednak to właśnie tutaj rozstrzyga się jedna z najważniejszych spraw współczesnej kultury: sposób, w jaki maszyny będą reprezentować sens. Jeśli sens zostaje zakodowany w geometrii, to pytanie o metrykę, model, indeks i korpus staje się pytaniem o warunki poznania. Brzmi to technicznie, lecz jest głęboko humanistyczne, gdyż ostatecznie chodzi o to, czy maszyna pomoże człowiekowi odnaleźć znaczenie, czy jedynie szybciej wyprodukuje jego pozór.

Borwankarowska wizja jest cenna, ponieważ prowadzi od podstaw matematyczno-programistycznych do praktycznych systemów i dalej – ku pytaniu o standard. Przedstawia Word2Vec i embeddings, ale nie zatrzymuje się na fascynacji arytmetyką pojęć. Wskazuje FAISS, lecz nie traktuje indeksu jako celu samego w sobie. Prezentuje SQLite-vss jako narzędzie osobistej

pamięci, a PostgreSQL z pgvector jako sposób budowy naukowego wyszukiwania wielopoziomowego.

Pokazuje RAG jako rygor źródłowy, a nie gadżet. Lokalne LLM traktuje jako decyzję o prywatności i kontroli. VQL widzi jako projekt wspólnej gramatyki, a nie gotowy mit założycielski. Z tej opowieści wyłania się model AI skromniejszej, lecz bardziej odpowiedzialnej postaci. Nie AI jako wszechwiedzącego kapłana chmury, lecz AI jako systemu pracy z własnym korpusem wiedzy. Nie generatora pewności, lecz asystenta potrafiącego wskazać źródła. Nie czarnej skrzynki, lecz łańcucha operacji: od dokumentu, przez embedding, wyszukiwanie i kontekst, aż po prompt, model, cytowanie i audyt. Nie jednego scentralizowanego megafonu, lecz pluralistycznej infrastruktury, którą można budować lokalnie, integrować hybrydowo i standaryzować tam, gdzie wymaga tego przenośność.

Wizja ta jest szczególnie istotna dziś, gdy generatywna AI często sprzedawana jest w formie spektaklu. Spektakl kocha natychmiastowość, wielkie obietnice, efektowne demonstracje i język przełomu. Jednak wiedza nie żyje ze spektaklu; żyje z pamięci, metody, źródeł, korekty, powtarzalności i sporu. Bazy wektorowe mogą pomóc wiedzy przetrwać w epoce nadmiaru informacji, o ile nie zostaną sprowadzone do kolejnego mechanizmu automatycznego streszczania treści dla ludzi, którzy nie chcą już czytać niczego. To byłaby cywilizacja skrótu, nie rozumu.

Dlatego najlepszym hasłem tej architektury mogłoby być: nie pytaj tylko modelu, pytaj pamięci. Model może mówić, ale pamięć może go ograniczyć. Źródło może go skorygować, metadane osadzić w czasie, a cytowanie – rozliczyć. Standard czyni proces przenośnym, lokalność chroni prywatność, a człowiek nadaje temu wszystkiemu sens. Dopiero suma tych elementów tworzy system zasługujący na poważne traktowanie.

Na poziomie praktycznym wnioski są jasne. Kto buduje systemy AI dla organizacji, powinien zacząć nie od wyboru najmodniejszego modelu, lecz od audytu korpusu wiedzy. Jakie dokumenty istnieją? Które są aktualne? Kto ma do nich dostęp? Jak są opisane i dzielone? Jakie metadane trzeba zachować? Jak sprawdzać retrieval i wymuszać cytowania? Jak mierzyć jakość odpowiedzi oraz postępować z dokumentami poufnymi? Jak odróżniać wersje robocze od finalnych, aktualizować embeddings czy wycofywać stare źródła? Jak szkolić użytkowników? Dopiero po udzieleniu odpowiedzi na te pytania wybór modelu ma sens. W przeciwnym razie zaczynamy budowę domu od wyboru dzwonka do drzwi.

AI jako rygorystyczny aparat wspomagania rozumu

W wymiarze naukowym wnioski są jednoznaczne. Systemy semantyczne mogą przyspieszyć przegląd literatury, lecz nie zastąpią krytyki metodologicznej. Potrafią ujawniać analogie, ale nie dowodzić prawdziwości; odnajdywać podobne prace, lecz nie rozstrzygać o ich wartości. Wspierają interdyscyplinarność, jednak nie zwalniają badacza z obowiązku rozumienia pojęć. Dobra nauka wiele zyska na bazach wektorowych, podczas gdy zła zyska jedynie szybszy sposób produkcji pozorów. Różnica nie tkwi w samej technologii, lecz w standardach pracy badawczej.

Na poziomie społecznym kluczowa jest kwestia władzy nad infrastrukturą wiedzy. Scentralizowanie systemów pamięci AI uczyni użytkowników zależnymi od kilku dostawców. Rozwój lokalnych i otwartych rozwiązań przyniesie natomiast większą różnorodność praktyk, języków, korpusów i zastosowań. Wprowadzenie standardów pokrewnych VQL ułatwi migrację, audyt i porównywanie systemów; ich brak skieruje rynek ku zamkniętym dialektom. A dialekt dostawcy to często elegancka nazwa smyczy.

W wymiarze kulturowym stawką jest to, czy AI wzmocni, czy osłabi praktykę czytania. RAG może zachęcać do sięgania po źródła, jeśli odpowiedzi prowadzą bezpośrednio do dokumentów, ale może też od nich odciągać, gdy streszczenie zastąpi kontakt z tekstem. Wszystko zależy od projektu interfejsu i obyczaju użytkownika. Dobra AI powinna być przewodnikiem po źródłach, a nie zasłoną spuszczoną między człowiekiem a tekstem. Powinna wskazywać: tu jest fragment, tu kontekst, tu ograniczenie, tu sprzeczność, tu trzeba doczytać. Zła AI powie: „nie martw się, już wszystko za ciebie zrozumiałam”. I właśnie wtedy należy się zmartwić.

Można zatem stwierdzić, że bazy wektorowe wprowadzają do informatyki nową wersję starego problemu hermeneutycznego: jak odnajdywać sens w zbiorze tekstów. Dawniej zajmowali się tym bibliotekarze, komentatorzy, indeksy i tradycje interpretacyjne. Dziś dołączają do nich embeddings, indeksy HNSW, wyszukiwanie hybrydowe i RAG.

Narzędzia ewoluują, lecz problem pozostaje ten sam: jak nie zgubić prawdy w nadmiarze informacji. Technologia daje nowe możliwości, ale nie znosi odpowiedzialności. Wręcz przeciwnie – ponieważ pozwala działać szybciej, wymaga surowszej dyscypliny. Finałowa teza artykułu brzmi zatem następująco: wektorowa infrastruktura AI jest obietnicą lepszej pamięci tylko wtedy, gdy zostanie podporządkowana źródłowości, prywatności, audytowi i standardom. Bez tego stanie się kolejnym narzędziem produkcji płynnych złudzeń. Z nimi – może stać się fundamentem nowej, bardziej odpowiedzialnej pracy z wiedzą. Nie dlatego, że maszyna stanie się mędrce, lecz dlatego, że człowiek zyska lepsze narzędzie odnajdywania i porządkowania własnej pamięci.

To być może najrozsądniejsza wizja sztucznej inteligencji: nie sztuczna wyrocznia, cyfrowy suweren czy chmurowy kapłan odpowiedzi, lecz cierpliwy, rygorystyczny, lokalnie kontrolowalny i źródłowo rozliczalny aparat wspomagania rozumu. Mniej fajerwerków, więcej kręgosłupa. Mniej „uwierz mi”, więcej „sprawdź tutaj”. Mniej halucynacyjnego baroku, więcej architektury dowodu. Jeśli AI ma naprawdę pomagać myśleć, musi najpierw nauczyć się pokory wobec źródeł. Pokora, nawet w technologii, bywa początkiem mądrości.

Ostatecznie wybór między wygodną chmurą a rygorem lokalnej architektury jest pytaniem o granice naszej podmiotowości. Czy wolimy AI jako lśniąca fasadę, która mówi nam to, co chcemy usłyszeć, czy jako surowy aparat dowodowy, który zmusza nas do powrotu do źródeł? Być może prawdziwa mądrość w erze algorytmów nie polega na znalezieniu modelu, który wie wszystko, lecz na zbudowaniu systemu, który ma odwagę powiedzieć: „nie wiem, ale sprawdź tutaj”.

Inspiracje

[1]



"Vector Databases A Practical Introduction" Nitin Borwankar

O'Reilly Media | ISBN: 9781098177591

Nitin Borwankar to doświadczony analityk danych i specjalista ds. baz danych z ponad trzydziestoletnim doświadczeniem w rozwoju i wdrażaniu rozwiązań dla przedsiębiorstw. Jest ceniony za swój wkład w edukację w zakresie data science, propagowanie narzędzi open source oraz pracę nad programami nauczania uczenia maszynowego. Borwankar ma doświadczenie inżynierskie i od dawna angażuje się w społeczność technologiczną, angażując się w projekty open source i pisząc artykuły techniczne. Jest autorem kilku książek poświęconych nowym technologiom i architekturze oprogramowania, w tym praktycznych poradników dotyczących wektorowych baz danych i rozwoju biznesu internetowego. Mieszkając w rejonie Zatoki San Francisco, nadal działa na styku architektury danych, sztucznej inteligencji i rozwoju oprogramowania.

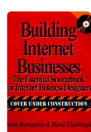
Nitin Borwankar is a seasoned data scientist and database professional with over three decades of experience in the development and implementation of enterprise data solutions. He is recognized for his contributions to data science education, his advocacy for open-source tools, and his work on machine learning curricula. Borwankar has a background in engineering and has been a long-time contributor to the tech community, including his involvement in open-source projects and technical writing. He is the author of several books focusing on emerging technologies and software architecture, including practical guides on vector databases and internet business development. Based in the San Francisco Bay Area, he continues to work at the intersection of data architecture, artificial intelligence, and software development.

University of Southern California

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] **"Building Successful Internet Businesses"**

David Elderbrock, Nitin Borwankar

1996 | Wiley Publishing | ISBN: 9780764570018

[2] "Reflections on Snowflakes"

Nitin Borwankar

1986

Literatura pokrewna (Google Scholar):

[3] "Building Complex Multi-Agent Systems Using"

Tim OBrien

[4] "Prompt Engineering for Llms Albert Ziegler"

[5] "Data Engineering with Generative n Agentic AI"

Justin J Leto

[6] "Agentic Mesh The GenAI-Powered Agent"

Eric Broda

[7] "Natural Language Processing in Action 2E Hobson Lane"

Artykuły naukowe

Artykuły pokrewne (Google Scholar):

[1] "The Laggard's Lock: A Theoretical Model of New Technology Adoption Constrained by Outdated Paradigms"

Alby Anand Kurian, Nitin Borwankar

2025 | SSRN Electronic Journal | DOI: 10.2139/ssrn.5398026

[Google Scholar](#)

[2] "A Study on Mahila Dakshata/Suraksha Samities of Maharashtra"

Meeran Borwankar

2012 | Global Community Policing | DOI: 10.1201/b12359-6

[Google Scholar](#)

[3] "Building Trust and Collaboration: Community Policing Initiatives"

Meeran Chadha Borwankar

2025 | Rethinking the Police for a Better Future | DOI: 10.1007/978-3-031-83173-7_6

[Google Scholar](#)

[4] "Viscoelastic properties of foods"

R.P. Borwankar

1993 | Trends in Food Science & Technology | DOI: 10.1016/0924-2244(93)90109-n

[Google Scholar](#)

[5] "A Novel Compact Convolutional Neural Network for Real-Time Nondestructive Evaluation of Metallic Surfaces"

Raunak Borwankar, Reinhold Ludwig

2020 | IEEE Transactions on Instrumentation and Measurement | DOI: 10.1109/tim.2020.2990541

[Google Scholar](#)

[6] "Tailoring Failure Interactions for Validation Testing of Composite Laminate Progressive Failure Models"

Pranav Borwankar, Satchi Venkataraman

2025 | AIAA Journal | DOI: 10.2514/1.j063540

[Google Scholar](#)

[7] "Food texture and rheology: A tutorial review"

Rajendra P. Borwankar

1992 | Journal of Food Engineering | DOI: 10.1016/0260-8774(92)90016-y

[Google Scholar](#)

[8] "Elections and Toxicity on Twitter"

Indulekha Guha, Sameer Borwankar

2024 | SSRN Electronic Journal | DOI: 10.2139/ssrn.4802794

[Google Scholar](#)

[9] "Foreword"

S.B. Borwankar, S.B. Borwankar

2012 | Succeed Or Sink | DOI: 10.1016/b978-1-84334-634-0.50012-3

[Google Scholar](#)

[10] "Meta's Community Notes program is promising, but needs to prioritize transparency"

Sameer Borwankar

2025 | DOI: 10.64628/aam.4scena5rf

[Google Scholar](#)

[11] "NDTNet: Optical Nondestructive Evaluation with Compact Convolutional Neural Network"

Raunak Borwankar, Reinhold Ludwig

2020 | 2020 IEEE International Instrumentation and Measurement Technology Conference (I2MTC) |

DOI: 10.1109/i2mtc43012.2020.9128739

[Google Scholar](#)



Treść tworzy Fundacja Dobre Państwo.

Redaguje i publikuje APA ONE, autorski system redakcyjny Fundacji oparty na AI.

APA ONE Publishing System

Fundacja Dobre Państwo © 2026

Article ID: e71601c9-5364-4a3a-8ab4-c7a90eodfd56