

Inżynieria systemów GenAI: od mitu agentowości do architektury komponentowej w oparciu o książkę Tima O'Briena Building Complex Multi-Agent Systems Using Pattern Prompting



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje przejście od marketingowego mitu „autonomicznych agentów” do rygorystycznego podejścia inżynierskiego w systemach GenAI. Autor argumentuje, że modele LLM nie są niezależnym intelektem, lecz niedeterministycznymi komponentami systemowymi (podobnie jak niestabilny API), które wymagają ścisłej kontroli, walidacji i architektury opartej na wzorcach. Przywołując koncepcje Tima O'Briena oraz raporty MIT NANDA, tekst wskazuje na lukę kompetencyjną między pokazami demo a produkcją. Główną tezą jest konieczność traktowania GenAI jako testu stresowego dla architektury organizacji – technologia ta nie naprawia chaosu w danych i procesach, lecz go uwypukla. Rozwiązaniem jest zastosowanie Pattern Prompting i znanych wzorców inżynierii oprogramowania, aby stworzyć przewidywalne i bezpieczne systemy klasy enterprise.

The article analyzes the transition from the marketing myth of 'autonomous agents' to a rigorous engineering approach in GenAI systems. The author argues that LLM models are not independent intellects, but non-deterministic system components (similar to an unstable API) that require strict control, validation, and pattern-based architecture. Citing Tim O'Brien's concepts and MIT NANDA reports, the text points to the competency gap between demo showcases and production. The main thesis is the necessity of treating GenAI as a stress test for organizational architecture—this technology does not fix chaos in data and processes, but rather

highlights it. The solution is the application of Pattern Prompting and known software engineering patterns to create predictable and secure enterprise-grade systems.

Współczesna fascynacja generatywną sztuczną inteligencją często opiera się na niebezpiecznym micie „autonomicznego intelektu”, który w rzeczywistości jest jedynie sprawnością retoryczną pozbawioną rozumienia świata. Autor tekstu stawia tezę, że aby GenAI przestała być źródłem efektownych, lecz ryzykownych demonstracji, a stała się bezpiecznym narzędziem produkcyjnym, musi zostać zdemitologizowana i potraktowana jako niedeterministyczny, zawodny komponent systemowy. Wymaga to przejścia od intuicyjnego „vibe codingu” do rygorystycznej dyscypliny inżynieryjnej, opartej na sprawdzonych wzorcach architektury oprogramowania (takich jak GoF czy EIP) oraz wysokiej jakości infrastrukturze danych. Kluczowe jest uznanie, że LLM nie posiada odpowiedzialności epistemicznej i nie może być źródłem prawdy; musi zatem zostać otoczony „rusztowaniem odpowiedzialności” – systemami walidacji, precyzyjnymi potokami ETL oraz architekturą agentową rozumianą jako mikroarchitektura komponentowa, a nie cyfrowy pracownik. Tylko poprzez zastąpienie antropomorficznych metafor twardym językiem inżynierii można zbudować systemy skalowalne, audytowalne i odporne na inherentne ryzyka technologii GenAI.

SŁOWA KLUCZOWE

inżynieria systemów GenAI

architektura komponentowa

Pattern Prompting

niedeterminizm LLM

Retrieval-Augmented Generation

RAG

Pattern-Guided Coding

wzorzec ReAct

szum semantyczny

dryf zależności

odpowiedzialność epistemiczna

idempotencja w AI

mikroarchitektura agenta

Topologos

walidacja danych z LLM

GenAI to ryzykowny komponent systemowy, a nie autonomiczny intelekt

GenAI między mitem a inżynierią. Dlaczego sztuczna inteligencja nie znosi braku dyscypliny

Generatywna sztuczna inteligencja weszła do języka biznesu tak gwałtownie, jak niegdyś chmura, blockchain czy automatyzacja. Każda epoka technologiczna niesie obietnicę zbawienia i swoje słowo-klucz, którym można na chwilę zastąpić strategię. Dziś takim słowem jest „agent”. Wczoraj każdy system miał być „smart”, potem „cloud-native”, następnie „blockchainowy” – teraz wszystko ma być „agentowe”.

To nie jest dowód postępu. Czasem to tylko dowód, że marketing nauczył się szybciej biegać niż architektura.

Punkt wyjścia jest jasny: GenAI ma ogromny potencjał, lecz przesłania go rynkowy szum, przesada i niedojrzałość wdrożeniowa. Kluczowym rozpoznaniem jest fakt, że problemem nie jest wyłącznie jakość samych modeli, lecz luka kompetencyjna między specjalistami od AI a inżynierami budującymi systemy produkcyjne. Modele imponują w demonstracjach, ale demo to nie produkcja. Demo ma zachwycić; system ma działać w poniedziałek rano – po aktualizacji API, przy wzroście ruchu, pod presją użytkownika, audytu, błędu sieciowego i rachunku za tokeny. Ta różnica między pokazem a systemem stanowi jeden z najważniejszych tematów współczesnej debaty o sztucznej inteligencji.

Raport MIT NANDA „The GenAI Divide: State of AI in Business 2025” wskazuje, że mimo ogromnych inwestycji tylko niewielka część pilotaży przynosi mierzalny wpływ na rachunek wyników. Opisano to jako podział między projektami zintegrowanymi z przepływami pracy a większością inicjatyw pozbawionych uchwytne go efektu biznesowego. Nie należy jednak czytać tej statystyki jako aktu oskarżenia przeciw AI. Bardziej precyzyjna interpretacja brzmi: zawodzi nie tyle model, ile sposób jego wprowadzenia do organizacji. Technologia nie zastępuje architektury – ona tylko bezlitośnie ujawnia jej brak.

GenAI jest lustrem instytucji. Organizacja z nieuporządkowanymi danymi, niejasnymi procesami i magicznym myśleniem menedżerskim nie stanie się nagle inteligentna dzięki podłączeniu modelu językowego. Przeciwnie: model powiększy jej chaos i opakuje go w piękne zdania. Brudne dane nie stają się czyste od samego przejścia przez API, a niejasny proces nie zmienia się w strategię tylko dlatego, że nazwiemy go „workflow agentowym”. Błąd organizacyjny nie znika, gdy zostanie wypowiedziany przez model z pewnością profesora i lekkością copywritera. Dlatego główna teza jest inżynierska, choć ma konsekwencje kulturowe: GenAI nie jest wyjątkiem od zasad budowy systemów, lecz ich testem stresowym. Jeśli systemy informatyczne wymagają kontroli wersji, obserwowalności, testowania, walidacji i bezpieczeństwa, to systemy GenAI wymagają tego wszystkiego jeszcze bardziej.

Nie dlatego, że są bardziej „inteligentne”, ale dlatego, że są mniej przewidywalne. Im bardziej komponent przypomina rozmówcę, tym łatwiej zapomnieć, że pozostaje komponentem – a to zapomnienie jest początkiem wielu katastrof wdrożeniowych. Właśnie tutaj pojawia się najważniejsze przesunięcie pojęciowe: LLM powinien być traktowany nie jako autonomiczny intelekt, lecz jako "źle zachowujący się punkt końcowy RESTful". To określenie celowo odczarowuje technologię. Odbiera modelowi aurę cyfrowego demiurga i przywraca go do świata architektury systemów rozproszonych.

Model jest usługą – kosztowną, zmienną, niedeterministyczną i podatną na opóźnienia. Może zwrócić odpowiedź błędną, niepełną, przekonująco fałszywą lub formalnie niezgodną ze schematem. W klasycznym ujęciu API, które raz zwraca poprawny JSON, raz poezję, raz halucynację, a raz odmawia odpowiedzi przez filtr bezpieczeństwa, zostałoby uznane za komponent wysokiego ryzyka.

Od magicznego intelektu do komponentu systemowego

W świecie GenAI ten sam fakt bywa sprzedawany jako "elastyczność", "kreatywność" czy "rozumowanie". Słowa bywają tu grzeczną maską dla niestabilności. Właśnie dlatego potrzebujemy języka wzorców inżynierskich – nie po to, by zdusić innowację proceduralną nudą, lecz by odróżnić rzeczywistą innowację od zwykłej improwizacji. Tim O'Brien, którego aparat pojęciowy stanowi jeden z fundamentów tej notatki, proponuje konkretną demistyfikację: budowę systemów GenAI w oparciu o znane wzorce inżynierii oprogramowania, zamiast polegania na gwałtownie zmieniających się frameworkach i efemerycznych metaforach agentowości. Książka "Building Complex Multi-Agent Systems Using Pattern Prompting" jest w tym kontekście przewodnikiem po tworzeniu bezpiecznych i skalowalnych systemów agentowych z użyciem sprawdzonych narzędzi i wzorców inżynierskich.

Nie chodzi o negację nowości GenAI, lecz o wprowadzenie jej do świata, w którym awaria ma konkretną przyczynę, komunikat posiada schemat, prompt ma właściciela, a decyzja architektoniczna zostawia trwały ślad. Kluczowym elementem tej propozycji jest stworzenie mostu terminologicznego. "Agent" przestaje być magiczną istotą podejmującą decyzje, a staje się komponentem. "Pamięć krótkotrwała" przestaje brzmieć jak cecha psychologiczna i zaczyna oznaczać stan sesji. "Pamięć długotrwała" nie jest już quasi-świadomością systemu, lecz systemem ewidencji. "Narzędzie" okazuje się adapterem, a "system wieloagentowy" – architekturą komponentową.

Ten zabieg, pozornie jedynie językowy, ma w istocie znaczenie dyscyplinujące. Kto zmienia słownik, zmienia sposób projektowania; kto nazywa zjawisko precyzyjnie, trudniej go uwieść bajką. W debacie o AI szczególnie niebezpieczne jest antropomorfizowanie. Gdy mówimy, że model "wie", "rozumie", "chce", "uznał", "zdecydował" lub "pomyślał", importujemy do inżynierii kategorie psychologiczne, które nadają systemowi pozór sprawstwa. W codziennym języku jest to wygodne, w architekturze produkcyjnej – zgubne. Model niczego nie "wie" w sensie odpowiedzialnego posiadania wiedzy; generuje odpowiedź na podstawie relacji statystycznych, parametrów, kontekstu i instrukcji. To wystarczy do tworzenia znakomitych streszczeń, klasyfikacji czy interfejsów językowych, ale nie pozwala na bezkontrolne zapisywanie danych w systemach ewidencji czy podejmowanie decyzji prawnych, medycznych i finansowych. Dlatego nad każdym projektem GenAI powinno wisieć ostrzeżenie, niczym tablica BHP nad wejściem do hali produkcyjnej: nie używaj danych wygenerowanych przez GenAI bezpośrednio w innych systemach IT. Każdy rekord z LLM musi przejść walidację przed zapisaniem w systemie ewidencji.

Konsekwencje tej zasady są daleko idące: LLM nie może być źródłem prawdy. Może służyć jako narzędzie transformacji, asystent interpretacji czy generator propozycji, ale nie jako samodzielny depozytariusz faktu. W tym miejscu RAG (Retrieval-Augmented Generation) przestaje być modnym skrótem, a staje się koniecznością architektoniczną. Skoro model nie ma dostępu do prawdy obiektywnej, trzeba go związać z systemem ewidencji. Jeśli ma tendencję do halucynacji, należy ograniczyć jego swobodę poprzez dołączanie źródeł. Jeśli operuje na prawdopodobieństwie językowym, musi otrzymać konkretne dane, które będą działały jak kotwice. RAG nie czyni modelu nieomylnym, ale sprawia, że przestaje być samotny. Samotny model językowy w środowisku produkcyjnym przypomina elokwentnego praktykanta zamkniętego w archiwum bez katalogu: mówi płynnie, ale niekoniecznie wie, na której półce leży dokument.

RAG nie rozwiązuje jednak wszystkiego. Czyste wyszukiwanie wektorowe zawodzi przy nazwach własnych, numerach seryjnych czy twardej logice zbiorów. Należy zatem myśleć o zintegrowanym wyszukiwaniu wiedzy: połączeniu semantyki wektorowej, słów kluczowych, metadanych i – tam, gdzie to uzasadnione – grafów wiedzy. Technologia uczy tu pokory: znaczenie nie jest jednowymiarowe, dokument nie jest workiem zdań, a wiedza organizacyjna czymś więcej niż chmurą podobieństw. Największym grzechem wielu wdrożeń jest redukcja świata do wektora i promptu. Wtedy projekt przypomina próbę zbudowania administracji państwowej z karteczek samoprzylepnych i dobrych intencji. Wektor świetnie reprezentuje podobieństwo semantyczne, ale nie zastąpi modelu danych, rejestru źródeł czy polityki dostępu. Embedding wskaże, że dwa fragmenty są blisko znaczeniowo, lecz nie powie, który dokument jest obowiązujący, a który został uchylony lub jest historycznym śladem błędu.

To prowadzi do tezy, że jakość GenAI jest funkcją jakości infrastruktury poznawczej organizacji. Jeśli dane są zduplikowane, sprzeczne i pozbawione metadanych, model będzie produkował jedynie elegancki chaos. To tzw. "semantic noise" – szum semantyczny. Brudne dane zatrują bazę wiedzy, a ta infekuje odpowiedzi modelu niezależnie od jego mocy. Nie istnieje model premium, który odkupiłby grzechy złego ETL-u.

LLM jako zawodny komponent systemowy zamiast intelektu

Tu kończy się dzieciństwo GenAI, a zaczyna inżynieria. Dzieciństwo polegało na pytaniu: „co model potrafi?”. Dorosłość zaczyna się od pytania: „w jakim systemie model może bezpiecznie pracować?”. To pytanie jest mniej widowiskowe, ale bardziej opłacalne. Nie brzmi jak hasło konferencyjne, lecz jak odpowiedzialność architekta. I właśnie dlatego jest ważniejsze.

Ontologia LLM: model językowy jako komponent ryzyka, a nie cyfrowy rozum. Jeśli pierwszym warunkiem dojrzałego wdrożenia GenAI jest wyjście z mgły marketingu, drugim jest prawidłowe nazwanie samego modelu. To pozornie akademicki gest, który w rzeczywistości ma znaczenie operacyjne. Źle nazwany komponent zostanie źle zaprojektowany, źle zabezpieczony i źle rozliczony. Jeśli nazwiemy LLM „partnerem poznawczym”, łatwo zaczniemy mu powierzać zadania wymagające sądu, rozumienia i odpowiedzialności. Jeśli jednak nazwiemy go „silnikiem predykcji językowej”, od razu zapytamy o walidację, źródła, błędy, granice, koszt, latencję i nadzór. W inżynierii słowa są tańsze niż awarie, ale potrafią być ich przyczyną.

Autor proponuje bardzo celną definicję roboczą: Large Language Model nie jest autonomicznym intelektem, lecz "źle zachowującym się punktem końcowym RESTful". To określenie może brzmieć brutalnie wobec technologii, która potrafi pisać sonety, streszczać umowy, generować kod i symulować rozmowę z erudytą. Ale jego siła polega właśnie na tej brutalności.

Dla architekta systemów pytanie nie brzmi: „czy model robi wrażenie?”, lecz: „jak zachowuje się jako zależność produkcyjna?”. A zachowuje się kłopotliwie: bywa niedeterministyczny, kosztowny, powolny, podatny na ograniczenia dostawcy, wrażliwy na konstrukcję promptu i zdolny do generowania błędnych odpowiedzi z pełną retoryczną pewnością. Ta ontologia jest kluczowa. Klasyczny punkt końcowy API powinien być przewidywalny: ten sam input w tych samych warunkach powinien zwrócić ten sam lub równoważny output; błędy powinny być typowane, odpowiedzi powinny respektować schemat, a latencja mieścić się w znanych granicach. Retry nie powinien generować nowych znaczeń. LLM narusza wiele z tych oczekiwań. Jego odpowiedź jest

funkcją promptu, kontekstu, parametrów generacji, wersji modelu i polityk bezpieczeństwa, a czasem również zmian wprowadzonych przez dostawcę bez pełnej widoczności dla użytkownika końcowego. Model może zostać „ulepszony” i tym samym zmienić zachowanie aplikacji, która jeszcze wczoraj przechodziła testy.

W klasycznej inżynierii nazwalibyśmy to ryzykiem dryfu zależności. W języku prezentacji sprzedażowych bywa to przedstawiane jako ciągły rozwój. Rozwój jest piękny, dopóki nie psuje produkcji. Dlatego podstawowa architektura systemów GenAI powinna zaczynać się od pokory wobec niedeterminizmu. Niedeterminizm nie jest wadą absolutną – jest źródłem elastyczności językowej, zdolności parafrazy, adaptacji stylu, generowania hipotez i pracy z nieustrukturyzowanym tekstem.

Problem zaczyna się wtedy, gdy niedeterminizm zostaje wprowadzony tam, gdzie organizacja potrzebuje determinizmu: w księgowości, rejestrach, decyzjach administracyjnych, klasyfikacji prawnej, rozliczeniach, systemach bezpieczeństwa, medycynie, finansach, ubezpieczeniach czy compliance. Tam model nie może być ostatnią instancją. Może być pomocnikiem, ale nie organem decyzyjnym. Może proponować, ale nie powinien sam wpisywać prawdy do księgi. Z tego powodu most terminologiczny zaproponowany w notatce nie jest ozdobą dydaktyczną, lecz narzędziem zarządzania ryzykiem. „Agent” staje się komponentem, „tool” adapterem, „pamięć krótkotrwała” stanem sesji, „pamięć długotrwała” systemem ewidencji, a „system wieloagentowy” architekturą komponentową. Ten przekład odbiera GenAI metafizyczną nadwyżkę. Nie mówi już: „agent pamięta i działa”, lecz: „komponent operuje na stanie sesji, wywołuje adaptery i korzysta z systemu ewidencji”. Różnica jest ogromna.

W pierwszym języku projektujemy cyfrowego pomocnika. W drugim – system rozproszony z zawodnym komponentem językowym. Właśnie dlatego „język antropomorficzny powinien być w dokumentacji technicznej traktowany jak substancja łatwopalna”. Można go używać potocznie, ale nie wolno budować na nim architektury.

Gdy mówimy, że model „pamięta”, często mamy na myśli, że w oknie kontekstowym znajduje się wcześniejszy fragment rozmowy albo że system dokleił do promptu dane z bazy. To nie jest pamięć w sensie psychologicznym, lecz mechanizm dostępu do kontekstu. Gdy mówimy, że model „rozumie dokument”, często oznacza to jedynie generowanie trafnego przekształcenia tekstu. To nie jest rozumienie semantyczne zakorzenione w świecie. Gdy mówimy, że model „wnioskuje”, często oznacza to układanie przekonującego łańcucha językowego. To nie musi być dowód – czasem jest to tylko elegancko ubrana asocjacja.

Inżynieria systemów wymaga zatem rozdzielenia trzech poziomów: zdolności językowej, zdolności operacyjnej i odpowiedzialności epistemicznej. LLM ma imponującą zdolność językową. Może streszczać, przekształcać, klasyfikować, tłumaczyć, generować warianty, symulować ton, układać argumentację i wspierać eksplorację materiału. Zdolność operacyjną uzyskuje dopiero wtedy, gdy zostanie podłączony do narzędzi, baz danych, kolejek, adapterów i procedur.

LLM jako ryzykowny komponent wymagający rusztowania odpowiedzialności

Model nie posiada odpowiedzialności epistemicznej sam z siebie. Może wskazać prawdopodobną odpowiedź, ale nie gwarantuje jej prawdziwości. Brzmi jak ekspert, lecz nie ponosi odpowiedzialności eksperta. Generuje wyjaśnienia, ale nie ma interesu w prawdzie poza tym, co wymusza architektura, instrukcje i walidacja. LLM jest zatem komponentem paradoksalnym: im lepiej mówi, tym większe stwarza ryzyko nadużycia zaufania.

Zwykły błąd w systemie informatycznym często objawia się jawnie: wyjątkiem, kodem błędu, pustą odpowiedzią lub awarią połączenia. Błąd LLM bywa aksamitny. Ma składnię, styl, pewność i pozór kompetencji. Może być fałszywy, nie będąc jednocześnie chaotycznym; może być niepoprawny, lecz piękny. A piękny błąd jest groźniejszy od brzydkiego, bo prześlizguje się przez uwagę człowieka jak dyplomata przez bankiet. Dlatego projektowanie systemów GenAI wymaga szczególnego rodzaju podejrzliwości. Nie chodzi o antytechnologiczną nieufność, lecz o profesjonalną sceptyczność. Architekt powinien pytać: co się stanie, gdy model odpowie niezgodnie ze schematem? Gdy poda wartość zmyśloną? Gdy odmówi odpowiedzi lub prompt injection zmieni jego zachowanie? Co w przypadku podania danych osobowych przez użytkownika? A jeśli model wygeneruje poprawną składniowo, lecz fałszywą klasyfikację? Co gdy retry powtórzy operację i utworzy drugi rekord, albo koszt pętli agentowej wymknie się spod kontroli? W tych pytaniach nie ma defetyzmu. Jest dorosłość. LLM bowiem nie gwarantuje czterech klasycznych fundamentów IT: szybkości, niezawodności, bezpieczeństwa i skalowalności.

Każdy z tych braków niesie osobną konsekwencję architektoniczną. Brak gwarantowanej szybkości wymusza asynchroniczność, kolejki i zarządzanie latencją. Brak niezawodności narzuca retry, circuit breaker, dead-letter queues i walidację odpowiedzi. Deficyt bezpieczeństwa wymaga izolacji danych, filtrowania wejścia i wyjścia, ochrony przed prompt injection oraz kontroli PII. Brak skalowalności wymusza budżetowanie tokenów, limity, dobór modeli do zadań i testy obciążeniowe. LLM nie jest więc centrum systemu, lecz komponentem, wokół którego trzeba zbudować rusztowanie odpowiedzialności.

Niektórzy entuzjaści powiedzą, że takie podejście odbiera AI rewolucyjność. Jest odwrotnie. Dopiero ono pozwala rewolucji przeżyć kontakt z rzeczywistością. Historia technologii zna wiele wynalazków, które stawały się społecznie przełomowe nie w chwili pierwszego zachwytu, lecz dopiero wtedy, gdy zostały ujęte w standardy, procedury i instytucje. Elektryczność potrzebowała sieci, norm bezpieczeństwa i liczników. Lotnictwo – kontroli ruchu, procedur serwisowych i czarnych skrzynek. Bankowość elektroniczna – kryptografii, audytu i regulacji. GenAI potrzebuje analogicznej dojrzałości: obserwowalności, walidacji, wersjonowania, testowania, polityk dostępu i odpowiedzialności za każdy etap przepływu informacji. Problem polega na tym, że kultura technologiczna ostatnich lat często premiuje demonstrację ponad utrzymanie. Łatwo pokazać agenta, który w kontrolowanych warunkach wykona zadanie; trudniej utrzymać system działający przez miesiące stabilnie, rozliczalnie i bezpiecznie. Łatwo przygotować film, na którym model rezerwuje podróż, pisze maila i analizuje raport. Trudniej odpowiedzieć, kto odpowiada za błąd w treści, kto autoryzuje płatność, gdzie zapisano dane i jak cofnięto błędną decyzję; czy użytkownik wiedział, że wynik jest probabilistyczny, i czy organizacja potrafi odtworzyć pełną ścieżkę decyzji.

Agent AI jako mikroarchitektura a nie cyfrowy pracownik

Demo kocha efekt. Produkcja kocha ślad. W tym punkcie kluczowe staje się precyzyjne rozumienie pojęcia „agenta”. W popularnym dyskursie agent AI bywa przedstawiany jako cyfrowy pracownik: otrzymuje cel, planuje, korzysta z narzędzi, obserwuje skutki i koryguje działanie. Z perspektywy inżynierskiej agent to jednak mikroarchitektura: pętla decyzyjna połączona z narzędziami przez adaptery, która zarządza stanem sesji i komunikuje się z otoczeniem według określonych wzorców.

Wzorzec ReAct opisuje tę strukturę jako pętlę „Myśl – Akcja – Obserwacja”, w której model analizuje kontekst, wywołuje narzędzie, otrzymuje wynik i aktualizuje stan. To użyteczna abstrakcja, pod warunkiem że nie pomylimy jej z psychologią. „Myśl” w ReAct nie jest procesem myślowym człowieka, lecz etapem generowania decyzji tekstowej lub strukturalnej. „Obserwacja” nie jest percepcją świata, a wynikiem wywołania narzędzia. „Akcja” nie wynika z woli, lecz jest komunikatem skierowanym do adaptera.

Takie doprecyzowanie nie pomniejsza wartości systemu – przeciwnie, umożliwia jego budowę. Jeśli „akcja” jest komunikatem, można ją walidować. Jeśli „obserwacja” jest wynikiem narzędzia, można ją typować. Jeśli „myśl” jest etapem decyzyjnym, można ją ograniczać schematem, temperaturą, krytykiem, regułami bezpieczeństwa czy mechanizmem human-in-the-loop. Gdy stan sesji traktujemy jako stan sesji, możemy określić jego czas życia, politykę czyszczenia i zakres danych.

Jeśli pamięć długotrwała jest systemem ewidencji, można przypisać jej właściciela i procedury audytu. Antropomorfizacja rozmywa odpowiedzialność; inżynieria ją lokalizuje.

Szczególnym problemem LLM jest brak idempotencji. W tradycyjnych systemach idempotencja oznacza, że powtórzenie tej samej operacji nie zmieni rezultatu ponad pierwszy efekt, co jest kluczowe przy retry, awariach sieciowych i przetwarzaniu komunikatów. LLM nie daje takiego komfortu. Powtórzenie promptu może przynieść inną odpowiedź, a ponowne uruchomienie agenta może wywołać inne narzędzie, inną klasyfikację lub inną decyzję. Jeśli taki model zostanie bezpośrednio podłączony do systemu wykonawczego, retry staje się ruletką. Dlatego system powinien oddzielać generowanie propozycji od zatwierdzania skutku: model przygotowuje kandydatów, a system decyduje, który z nich jest poprawny i bezpieczny.

Z tym wiąże się rola temperatury jako parametru inżynierskiego. W potocznym ujęciu temperatura bywa sprowadzana do „kreatywności”, jednak w architekturze systemów agentowych służy kontroli wariancji. Ruter powinien pracować z temperaturą bliską zeru, aby generować przewidywalne decyzje; krytyk lub sędzia wymaga niskiej temperatury dla stabilności oceny; edytor może korzystać z wyższej wartości, jeśli celem jest różnorodność językowa. Nie istnieje jedna „najlepsza” konfiguracja modelu – istnieje konfiguracja właściwa dla funkcji komponentu. Kreatywność w ruterze jest błędem, natomiast powtarzalność w generatorze wariantów stylistycznych może być wadą.

Architektura zaczyna się tam, gdzie przestajemy używać jednego młotka do wszystkich śrub. Podobnie należy podchodzić do doboru modeli. Największy model nie zawsze jest najlepszym; często bywa najdroższym sposobem wykonania prostego zadania. Do klasyfikacji intencji, ekstrakcji pól czy routingu zazwyczaj wystarczy mniejszy i szybszy model. Większe modele należy rezerwować dla złożonej syntezy, analizy wielokrokowej lub pracy z niejednoznacznym materiałem. Ekonomia GenAI jest zatem integralną częścią architektury, a nie dodatkiem po wdrożeniu. Źle zaprojektowany system agentowy może nie tylko halucynować, ale i przepalać budżet z gorliwością koparki kryptowalut w piwnicy entuzjasty.

W praktyce oznacza to konieczność otoczenia LLM warstwami kontroli. Na wejściu: filtry, klasyfikacja danych, ochrona przed prompt injection i anonimizacja. Wewnątrz: ograniczenia kontekstu, instrukcje systemowe, wersjonowane prompty i kontrola parametrów. Na wyjściu: walidatory strukturalne, testy zgodności, krytyk oraz – w zastosowaniach krytycznych – decyzja człowieka. Taki system nie jest mniej „AI”, lecz bardziej odpowiedzialny. Sztuczna inteligencja bez architektury przypomina silnik odrzutowy przykręcony do drewnianego wozu. Robi wrażenie przez kilka sekund.

Ta perspektywa przesuwca ciężar kompetencji. W pierwszej fazie fascynacji wydawało się, że kluczowy będzie „prompt engineering” jako sztuka rozmowy z modelem. W fazie produkcyjnej okazuje się, że ważniejsza jest inżynieria systemowa: projektowanie przepływów, danych, kolejek, walidacji, bezpieczeństwa i obserwowalności.

Przejście od magii promptowania do dyscypliny inżynierskiej

Prompt pozostaje ważny, lecz przestaje być zaklęciem. Staje się artefaktem konfiguracyjnym, częścią kodu — elementem wersjonowanym i testowanym. Magia zostaje zdegradowana do repozytorium. I bardzo dobrze. W produkcji magia powinna mieć numer wersji.

Tu zaczyna się rola Pattern-Guided Coding, czyli kodowania opartego na wzorcach. Jeśli LLM jest komponentem ryzyka, a systemy agentowe mikroarchitekturami, potrzebujemy języka pozwalającego projektować je w sposób powtarzalny. Wzorce GoF oraz Enterprise Integration Patterns pełnią tu funkcję wspólnej gramatyki. Pozwalają precyzyjnie określić: tutaj potrzebujemy adaptera, tutaj strategii, tutaj command, tutaj circuit breaker, tutaj competing consumers, tutaj dead-letter channel, tutaj proxy bezpieczeństwa, a tutaj template method dla konsumentów. Wzorce nie są modą; są pamięcią inżynierii. Pamięcią zrodzoną z awarii, które ktoś już przeżył, opisał i przekuł w rozwiązanie.

Największą zaletą takiego podejścia jest jawność decyzji projektowych. Swobodny "vibe coding" może wygenerować kod działający w przykładzie, ale nie ujawnia przyjętych założeń architektonicznych. Pattern-Guided Coding wymusza nazwanie decyzji przed implementacją. Jeśli używamy kolejki błędów, musimy określić jakiej i po co. Przy stosowaniu retry trzeba zdefiniować limity i backoff. Wprowadzając krytyka, należy ustalić, co weryfikuje i jakie są kryteria odrzucenia. Gdy model ma dostęp do narzędzi, konieczne jest określenie adapterów, uprawnień i skutków ubocznych. Projekt przestaje być rozmową z czarną skrzynką, a staje się audytowalnym procesem.

Z perspektywy instytucjonalnej jest to kwestia odpowiedzialności. Organizacje nie mogą wdrażać GenAI jako "czyjejś zabawki z działu innowacji", która po udanym pokazie zostaje przerzucona do produkcji. Każdy element powinien mieć właściciela: prompt, model, baza wektorowa, pipeline ETL, polityka bezpieczeństwa, walidator, monitoring kosztów, zestaw testów oraz procedura rollbacku.

Odpowiedzialność rozproszona jest często odpowiedzialnością niczyją, a w systemach GenAI szybko staje się ona cudzą szkodą. Ontologia LLM prowadzi nas zatem do etyki inżynierskiej. Nie

chodzi o wielkie deklaracje dotyczące przyszłości ludzkości, lecz o etykę codziennego projektowania: nie zapisuj niezweryfikowanego wyniku do systemu ewidencji; nie wysyłaj danych osobowych do modelu bez podstawy i zabezpieczeń; nie ukrywaj niedeterminizmu przed użytkownikiem; nie udawaj pewności, gdy masz tylko prawdopodobieństwo; nie pozwalaj agentowi działać poza zakresem uprawnień; nie myl demonstracji z dowodem produkcyjnej gotowości. To małe zasady, ale z nich składa się duże zaufanie.

LLM jest narzędziem o wysokiej sprawności retorycznej i ograniczonej odpowiedzialności poznawczej. Potrafi operować językiem z efektywnością, której nie sposób zignorować, jednak nie posiada naturalnego kontaktu ze światem, nie gwarantuje prawdy, nie rozumie społecznych konsekwencji swoich odpowiedzi i nie odróżnia błędu od faktu w sposób wymagany przez systemy krytyczne. Właśnie dlatego nie należy pytać, czy LLM jest "inteligentny" w sensie publicystycznym.

To pytanie częściej zaciemnia, niż wyjaśnia. Należy pytać o kontrakty wejścia i wyjścia, generowane błędy, sposób walidacji, ograniczanie skutków ubocznych, monitorowanie dryfu, ochronę danych oraz o to, kto odpowiada za decyzję wykonaną przez system. Dopiero tak zdefiniowany model może stać się użytecznym elementem infrastruktury. Bez tego staje się obiektem kultu — a obiekty kultu mają jedną wadę: rzadko dobrze przechodzą testy obciążeniowe.

Płynność językowa to nie rozumienie świata

Należy teraz przejść od ontologii inżynierskiej do granic poznawczych. Skoro LLM jest komponentem językowym, warto zapytać, czego ta językowość nie obejmuje: czym różni się syntaktyczna płynność od semantycznego rozumienia, dlaczego problem ugruntowania symboli nie jest jedynie akademicką ciekawostką i skąd biorą się halucynacje. To pozwala zrozumieć, dlaczego „rozumowanie” modelu wymaga opisu znacznie ostrożniejszego niż ten, którym karmi nas marketing sztucznej inteligencji.

Płynność języka nie jest tożsama z rozumieniem świata. Traktując model jako komponent systemowy, a nie autonomiczny intelekt, musimy precyzyjnie określić jego granice poznawcze. Nie chodzi tu o publicystyczny spór o to, czy sztuczna inteligencja „myśli” — pytanie to zazwyczaj więcej zdradza o pragnieniach pytającego niż o samej technologii. Chodzi o kwestie bardziej konkretne: jakie operacje wykonuje LLM, czego te operacje nie gwarantują i dlaczego różnica między generowaniem poprawnych sekwencji językowych a rozumieniem znaczenia ma bezpośrednie konsekwencje dla projektowania systemów produkcyjnych.

LLM są mechanizmami statystycznymi, które generują płynność językową bez sprawdzalnego zakotwiczenia w treści. Modele przewidują kolejny token na podstawie rozkładów prawdopodobieństwa i kontekstu, a nie dlatego, że posiadają intencję, doświadczenie świata, świadomość znaczenia czy odpowiedzialność za prawdę. Takie rozpoznanie nie jest obelgą wobec technologii, lecz warunkiem jej poprawnego użycia.

Termometr nie jest gorszy od lekarza tylko dlatego, że nie rozumie gorączki; jest użyteczny, gdy wiemy dokładnie, co mierzy i kiedy jego wynik wymaga interwencji specjalisty. W przypadku LLM problem polega na tym, że instrument pomiarowy mówi ludzkim językiem. A kiedy coś przemawia w ten sposób, niemal automatycznie przypisujemy temu czemuś intencję i podmiotowość. Jest to nasza stara antropologiczna skłonność: słyszymy głos, więc szukamy osoby.

W systemach GenAI ta skłonność staje się ryzykiem projektowym. Model nie tylko zwraca wynik – on robi to w formie wypowiedzi często eleganckiej, pewnej siebie i retorycznie kompletnej. W efekcie błąd maszyny nie wygląda jak usterka techniczna, lecz jak opinia eksperta. A opinia eksperta, nawet fałszywa, ma w organizacji większą siłę perswazyjną niż komunikat „Error 503”.

Dlatego kluczowym pojęciem w krytyce LLM jest metafora „stochastycznej papugi”. Spopularyzowane w debacie naukowej przez Emily M. Bender, Timnit Gebru i współautorki określa to, że wielkie modele językowe mogą wytwarzać płynne formy bez gwarancji rozumienia czy zakotwiczenia w rzeczywistości. Model powtarza, przekształca i komponuje wzorce językowe, ale nie posiada semantycznego doświadczenia świata. Papuga jest tu oczywiście niesprawiedliwa wobec prawdziwych papug, które przynajmniej mają ciało, środowisko, głód, strach i gałąź pod pazurem.

Model ma jedynie wektory. Z tej różnicy bierze się problem ugruntowania symboli (symbol grounding problem). W ludzkim poznaniu słowa są powiązane z doświadczeniem cielesnym, społecznym i praktycznym. Słowo „woda” nie jest dla człowieka tylko elementem słownika; jest zimnem szklanki, ciężarem wiadra, zagrożeniem powodzią czy wspomnieniem dzieciństwa. Dla modelu językowego „woda” to przede wszystkim punkt w przestrzeni relacji między tokenami. Model może znakomicie pisać o wodzie, ale nie wie, czym jest mokość; może opisać ogień, lecz nie boi się poparzenia.

Syntaktyczna sprawność nie gwarantuje semantycznej wiarygodności

Może zdefiniować ciężar, ale niczego nie dźwiga. Ten brak ugruntowania nie jest filozoficzną ciekawostką z seminariów, na których kawa stygnie szybciej niż spór o świadomość; to konkretny problem operacyjny. Jeśli model nie ma kontaktu z fizycznością, instytucjonalnością i sprawczością świata, jego „rozumowanie” o skutkach działań może być przekonujące językowo, a zarazem kruche praktycznie. Model potrafi zasugerować procedurę brzmiącą logicznie, która jednak pomija ograniczenia prawne. Może zaproponować kod wyglądający poprawnie, lecz nieuwzględniający warunków brzegowych. Potrafi streścić dokument, gubiąc przy tym kluczowy wyjątek zapisany w jednym zdaniu. Może sklasyfikować przypadek, nie rozumiejąc realnego kosztu błędnej decyzji dla człowieka, firmy czy instytucji. Brak ugruntowania symboli prowadzi zatem do braku naturalnej odpowiedzialności za konsekwencje.

Warto tu przywołać eksperyment myślowy Johna Searle'a – chiński pokój. Jego sens jest prosty, choć przez dekady obudowano go labiryntem komentarzy: poprawna manipulacja symbolami nie jest tożsama z rozumieniem znaczenia. Osoba zamknięta w pokoju może, korzystając z instrukcji, przetwarzać chińskie znaki i zwracać odpowiedzi, które dla zewnętrznego obserwatora wyglądają na sensowne, mimo że sama nie zna języka chińskiego. Notatka źródłowa słusznie streszcza ten argument formułą: *syntax is not semantics*. W odniesieniu do LLM oznacza to, że poprawność zachowania językowego nie rozstrzyga o obecności rozumienia. System może przejść test płynności, nie przechodząc testu semantycznej odpowiedzialności. Oczywiście zwolennik silniejszych interpretacji AI odpowie, że człowiek również operuje symbolami, że znaczenie wyłania się z użycia i że nie mamy bezpośredniego dostępu do cudzego rozumienia – a zatem wystarczająco złożona manipulacja symbolami może być funkcjonalnie nieodróżnialna od rozumienia.

To istotny kontrargument, którego nie należy zbywać machnięciem ręki. W praktyce jednak systemy produkcyjne nie potrzebują metafizycznej deklaracji o tym, że model „naprawdę rozumie”. Potrzebują gwarancji działania. Jeśli system ma podejmować decyzje, klasyfikować dokumenty, odpowiadać klientom czy wspierać lekarza, prawnika i analityka, kluczowe jest nie to, czy jego zachowanie przypomina rozumienie, lecz czy wynik jest sprawdzalny, odtwarzalny, zgodny ze źródłami i bezpieczny. Samą płynność językową uznać za taką gwarancję nie można. Z tej perspektywy spór Searle'a nie musi kończyć się rozstrzygnięciem filozoficznym; wystarczy wyciągnąć z niego wniosek inżynierski: nie wolno mylić syntaktycznej sprawności z semantyczną wiarygodnością. Model może produkować odpowiedzi zgodne z gramatyką i stylem danej dziedziny, ale to nie oznacza, że wynik zakorzeniono we właściwym dokumencie, obowiązującym

stanie prawnym, aktualnym rejestrze czy empirycznym pomiarze. Dlatego architektura GenAI musi dostarczać to, czego model sam nie posiada: źródła, walidację, procedury, kontrolę wersji oraz pełną rozliczalność i możliwość zakwestionowania wyniku.

LLM to statystyczna emulacja, a nie rozumienie

Podobną funkcję pełni w notatce odwołanie do Noama Chomsky'ego, który od dekad krytycznie odnosi się do redukcji języka do statystyki zachowań. W ujęciu przywołanym w materiale źródłowym istnieje zasadnicza różnica między biologiczną zdolnością dziecka do nabywania mowy a maszynowym uczeniem opartym na ogromnych zbiorach danych. Dziecko nie potrzebuje połączyć internetu, by zacząć tworzyć zdania.

Maszyna natomiast wymaga kolosalnych korpusów, parametrów i energii obliczeniowej, by emulować powierzchowne aspekty kompetencji językowej. Nie chodzi tu o romantyczne wywyższanie człowieka, lecz o wskazanie różnicy w mechanizmach, co ma kluczowe znaczenie praktyczne. Ludzka kompetencja językowa jest osadzona w ciele, interakcji, intencjach, normach społecznych i wspólnym doświadczeniu. Model językowy uczy się jedynie z tekstowych śladów tych procesów. To tak, jakby rekonstruować życie miasta wyłącznie z paragonów, ogłoszeń, stenogramów i rozmów podsłuchanych przez ścianę. Można w ten sposób odkryć wiele wzorców i przewidzieć zachowania, ale nadal brakuje bezpośredniego uczestnictwa w tym, co tekst jedynie reprezentuje. LLM jest genialnym archeologiem dyskursu, ale nie jest obywatelem świata, który opisuje.

Innego rodzaju granica pojawia się przy odwołaniu do Gödla i Penrose'a. W notatce wskazano argument, według którego ludzkie myślenie może zawierać aspekty nieobliczalne lub nieredukowalne do formalnej procedury manipulacji symbolami. Jest to obszar bardziej spekulatywny, wymagający ostrożności. Nie trzeba jednak przyjmować całego programu Penrose'a, by uznać skromniejszą tezę: systemy obliczeniowe mają swoje granice, a modele statystyczne nie stają się podmiotami rozumienia tylko dlatego, że ich skala przekroczyła próg imponujący dla obserwatora. Wielkość modelu poprawia zakres i jakość predykcji, ale sama w sobie nie jest dowodem świadomości, intencjonalności ani niezawodnej dedukcji.

Wspólny rdzeń tych argumentów można ująć następująco: LLM posiada kompetencję dystrybucyjną, lecz brakuje mu pełnej kompetencji ontologicznej. Zna – w cudzysłowie – relacje między znakami i kontekstami użycia, ale nie zna świata jako przestrzeni oporu i konsekwencji. Potrafi przewidzieć prawdopodobny ciąg słów, lecz nie posiada naturalnego mechanizmu odróżniania prawdy od dobrze brzmiącej kontynuacji, o ile architektura mu tego nie narzuci.

Prawda w systemach GenAI nie jest więc immanentną cechą modelu, lecz wynikiem całego układu: danych, źródeł, retrievalu, promptu oraz kontroli i odpowiedzialności człowieka.

Tu pojawia się fenomen halucynacji. Nie jest ona przypadkowym defektem w rodzaju literówki, lecz konsekwencją sposobu działania modelu, który został wytrenowany do generowania prawdopodobnych odpowiedzi, a nie do milczenia w obliczu niewiedzy. Choć nowoczesne modele lepiej sygnalizują niepewność i odmawiają spekulacji, skłonność do produkowania płynnego wyniku pozostaje wpisana w ich naturę. Model często bardziej „chce” dokończyć wypowiedź, niż zatrzymać się przy granicy wiedzy. To „chce” jest oczywiście metaforą, lecz trafnie opisuje presję funkcji generatywnej: kontynuuj tekst.

Najgroźniejsze w halucynacjach LLM nie jest to, że model się myli – systemy informatyczne zawsze mogą popełniać błędy. Najgroźniejsze jest to, że model myli się narracyjnie. Produkuje błąd wraz z uzasadnieniem, kontekstem i pozorem źródłowości. Zwykła baza danych w przypadku braku rekordu zwróci pusty wynik; model językowy może w takiej sytuacji stworzyć rekord z mgły. To różnica między ciszą a konfabulacją. W świecie informacji cisza bywa irytująca, ale konfabulacja bywa kosztowna.

Dlatego słusznie łączy się halucynacje z brakiem skalibrowanej pewności. Model może wypowiadać błąd z taką samą pewnością językową jak fakt. Człowiek dysponuje sygnałami niepewności: waha się, sprawdza, odwołuje do doświadczenia. Model może zostać poinstruowany, by komunikować brak pewności, ale jest to zachowanie generowane, a nie wewnętrzny stan epistemiczny. Mówi „nie jestem pewien” tylko wtedy, gdy kontekst czyni taką odpowiedź prawdopodobną. W systemach krytycznych to za mało – niezbędna jest zewnętrzna kalibracja: porównanie ze źródłem, walidacja formalna i eskalacja do człowieka.

Z tym wiąże się luka rozumowania (The Reasoning Gap). Modele mogą sprawiać wrażenie logicznego myślenia, zwłaszcza gdy generują łańcuch argumentacji krok po kroku. Jednak taki ciąg nie zawsze jest dowodem w sensie logicznym; może być jedynie racjonalizacją wyniku powstałego heurystycznie, zawierającą ukryty błąd lub fałszywą analogię. Notatka wskazuje trzy obszary zawodności: systematyczne rozumowanie matematyczne, spójność logiczną w długim kontekście oraz generowanie nowatorskich wglądów. Matematyka jest tu najlepszym testem, gdyż nie daje się uwieść stylowi. W eseju słabość rozumowania można ukryć pod rytmem zdania; w rachunku wynik albo jest poprawny, albo nie.

Modele językowe znakomicie objaśniają znane procedury matematyczne, ale w zadaniach wymagających precyzyjnego dowodu często myślą heurystykę z algorytmem. Potrafią poprawnie

opisać metodę, by następnie błędnie ją wykonać lub wygenerować dowód, który brzmi przekonująco, lecz zawiera krytyczną lukę.

Kompensowanie braków modelu poprzez rygorystyczną architekturę systemową

To nie oznacza, że LLM nie nadają się do pomocy w matematyce czy programowaniu. Oznacza to raczej, że potrzebują narzędzi weryfikujących: interpreterów, solverów, systemów dowodowych, testów jednostkowych, kompilatorów i ludzi, którzy wiedzą, gdzie model może "ładnie zgubić minus". Drugim problemem jest spójność w długim kontekście. Im obszerniejszy materiał, tym trudniej utrzymać wszystkie zależności, wyjątki, definicje i ograniczenia. Model może poprawnie zidentyfikować główną tezę, ale przeoczyć warunek poboczny. Może na początku tekstu przyjąć jedno znaczenie pojęcia, by później nieświadomie przesunąć je w inne. Potrafi zintegrować dwa dokumenty, lecz jednocześnie znieść między nimi napięcie tam, gdzie należało je zachować. W długich projektach publikacyjnych, prawnych, analitycznych czy technicznych problemem nie jest samo to, "czy model umie pisać", lecz to, czy utrzyma architekturę znaczeń przez kilkadziesiąt stron. Bez planu, kontroli pojęć, cytowania źródeł i kolejnych rewizji nie ma na to gwarancji. Trzecia kwestia – nowatorskie wglądy – wymaga szczególnej ostrożności.

LLM potrafią tworzyć kombinacje, analogie, syntezy i hipotezy, które dla człowieka bywają inspirujące, wspierając tym samym proces twórczy. Ich nowość jest jednak statystycznie zależna od przetworzonych wzorców. Model może zaskakiwać, łącząc odległe rejestry językowe i domeny wiedzy, lecz nie oznacza to automatycznie, że generuje wiedzę w takim sensie, w jakim uczony odkrywa nowy mechanizm, prawnik tworzy przełomową konstrukcję, a inżynier rozwiązuje problem poprzez kontakt z oporem materiału. LLM jest znakomitym akceleratorem rekombinacji. Nie należy jednak mylić rekombinacji z odkryciem, choć czasem pierwsza może pomóc człowiekowi dojść do drugiego. Właśnie dlatego najrozsądniejsza koncepcja LLM nie jest ani pogardliwa, ani bałwochwalcza. Model nie jest głupią zabawką, ale nie jest też cyfrowym mędrce. To potężne narzędzie pracy na języku, ujawniające siłę statystycznego modelowania kultury tekstowej. Może być asystentem, tłumaczem, redaktorem, pomocnikiem programistycznym, klasyfikatorem, interfejsem do danych, generatorem wariantów czy przewodnikiem po dokumentach. Jego wartość rośnie jednak wtedy, gdy architektura kompensuje jego braki. LLM bez źródeł fantazjuje. LLM bez walidacji ryzykuje. LLM bez ograniczeń kosztuje. LLM bez nadzoru nabiera pewności siebie jak stażysta po pierwszym wystąpieniu na konferencji.

RAG jest właśnie odpowiedzią na część tych ograniczeń, przy czym jest to odpowiedź techniczna, a nie metafizyczna. Retrieval-Augmented Generation nie sprawia, że model zaczyna rozumieć świat; sprawia jedynie, że otrzymuje on fragmenty świata zapisane w kontrolowanym systemie wiedzy. Gdy użytkownik zadaje pytanie, system wyszukuje odpowiednie dokumenty, dołącza je do kontekstu i wymusza odpowiedź opartą na źródłach. Nie eliminuje to halucynacji całkowicie, ale zmniejsza ich prawdopodobieństwo i ułatwia kontrolę.

Model zostaje przywiązany do dokumentu jak latawiec do sznurka. Nadal może szarpać się na wietrze, ale trudniej mu odlecieć w czystą fantazję. Jednak RAG działa tylko wtedy, gdy sprawnie funkcjonuje retrieval. A ten z kolei zależy od tego, czy dane są dobrze przygotowane, oczyszczone, podzielone, opisane i wyszukiwane właściwą metodą. Jeśli baza wiedzy składa się z chaotycznych dokumentów, model będzie cytował chaos. Jeśli chunking przetnie sens na pół, model otrzyma półsens i wygeneruje całe zdanie. Jeśli embedding nie rozróżni nazw własnych, model może znaleźć dokument podobny, lecz niewłaściwy. Jeśli system pominie metadane, może przywołać dokument przestarzały zamiast aktualnego. RAG nie jest więc magicznym lekarstwem.

Jest dyscypliną zarządzania kontekstem. W tym miejscu filozofia poznania spotyka się z bardzo przyziemnym ETL-em. Problem symbol grounding mówi: model nie ma bezpośredniego kontaktu ze światem. ETL odpowiada: wobec tego trzeba starannie przygotować tekstowe reprezentacje świata. Problem halucynacji mówi: model generuje prawdopodobne kontynuacje. Walidacja odpowiada: wobec tego trzeba sprawdzić, czy kontynuacje są zgodne ze źródłami i schematem. Problem luki rozumowania mówi: model może mylić łańcuch słów z dowodem. Testy odpowiadają: wobec tego trzeba używać narzędzi formalnych, benchmarków, kompilatorów, solverów i ludzkiej recenzji. To jest najważniejszy sens inżynierii GenAI: nie kłócić się z naturą modelu, lecz zbudować system, który zna jego ograniczenia. Krytyka LLM nie musi prowadzić do pesymizmu; przeciwnie, rzetelna krytyka jest warunkiem produktywności. Kto wierzy w niemal wszechwiedzący model, szybko powierzy mu zadania, których ten nie powinien wykonywać, by rozczarować się przy pierwszej poważnej awarii. Kto rozumie jego ograniczenia, potrafi znaleźć dla niego właściwe miejsce. To trochę jak z bardzo utalentowanym, ale roztargnionym współpracownikiem: można z nim zdziałać wiele, o ile nie pozwoli mu się samodzielnie podpisywać umów, księgować faktur i przepisywać prawa budowlanego z pamięci.

LLM jako komponent syntezy a nie magazyn wiedzy

Granice poznawcze LLM prowadzą do praktycznej zasady: model powinien być używany tam, gdzie koszt błędu jest kontrolowalny, wynik można sprawdzić, a jego moc językowa daje przewagę nad

klasycznym oprogramowaniem. Nie warto stosować LLM do wszystkiego; byłaby to nie innowacja, lecz technologiczna nerwica natręctw. Jeśli problem da się rozwiązać prostą regułą, zapytaniem SQL, walidatorem, parserem lub klasycznym algorytmem, użycie modelu może zwiększyć koszt i ryzyko bez proporcjonalnej korzyści.

Modele językowe są najcenniejsze w obszarach nieustrukturyzowanego języka, wieloznaczności, syntezy, klasyfikacji semantycznej, transformacji stylu czy interfejsów naturalnych. Zasada ta ma także wymiar kulturowy. Współczesny hype wokół AI często opiera się na odwróceniu ciężaru dowodu: zakłada się, że skoro model potrafi wiele, należy go wdrażać wszędzie, a sceptyk ma tłumaczyć się z ostrożności. Dojrzała inżynieria powinna przywrócić właściwy porządek – to entuzjasta wdrożenia musi wykazać, że LLM jest odpowiednim narzędziem dla konkretnego zadania, że wynik podlega walidacji, dane są bezpieczne, a koszty policzone. Musi udowodnić, że błędy mają ścieżkę obsługi, użytkownik rozumie ograniczenia, a organizacja potrafi cofnąć szkodliwą zmianę. Innowacja bez ciężaru dowodu jest tylko kosztownym życzeniem. Walka z hype'em nie jest zatem walką przeciwko AI, lecz walką o AI nadające się do świata rzeczywistego. Sztuczna inteligencja nie potrzebuje kapłanów od cudów, lecz inżynierów od ograniczeń.

Nie potrzebuje deklaracji, że "zmieni wszystko", lecz odpowiedzi na pytania: co dokładnie zmienia, w jakim procesie, z jakim ryzykiem i kosztem, pod czyją odpowiedzialnością oraz z jaką procedurą awaryjną. Dopiero wtedy GenAI przestaje być widowiskiem, a staje się infrastrukturą. RAG jako inżynierska proteza prawdy: retrieval, chunking, embeddings i kontrolowane okno na świat.

Jeżeli LLM nie posiada naturalnego dostępu do świata, trzeba mu ten świat dostarczyć w formie kontrolowanej i sprawdzalnej. Choć brzmi to jak techniczna oczywistość, zdanie to zawiera istotę architektury RAG. Retrieval-Augmented Generation nie jest eleganckim dodatkiem ani dekoracją w prezentacji o "wiedzy firmowej". Jest próbą rozwiązania jednego z najważniejszych deficytów LLM: model mówi płynnie, lecz sam z siebie nie wie, na czym oprzeć prawdę. RAG przypina go zatem do źródła, dokumentu, systemu ewidencji czy konkretnego korpusu danych. Model przestaje być samotnym mówcą, a staje się komentatorem dostarczonego materiału.

Taka perspektywa ujmuje RAG jako konieczność, a nie opcję. W systemach produkcyjnych LLM nie powinien służyć jako wewnętrzny magazyn wiedzy, gdyż jego "wiedza" jest statystyczną destylacją danych treningowych – często nieaktualnych, niekompletnych i niedostępnych do audytu. RAG przekształca model z zawodnego generatora w komponent przypominający wzorec Content Enricher: system najpierw pobiera właściwe fakty z zewnętrznego źródła, a dopiero potem pozwala modelowi wykorzystać zdolności językowe do syntezy, objaśnienia lub transformacji.

RAG jako kotwica wiedzy i ryzyko błędnego chunkingu

Model bez retrievalu odpowiada z siebie – czerpie z mieszaniny danych treningowych, kontekstu rozmowy, instrukcji i prawdopodobieństwa językowego. Model wsparty RAG opiera się na materiale, który system wyszukał i dostarczył mu jako kontekst. Różnica przypomina tę między człowiekiem improwizującym z pamięci a kimś, kto ma przed sobą akta sprawy. Oczywiście posiadanie akt nie gwarantuje mądrości; można źle przeczytać dokument, pominąć wyjątek, zacytować niewłaściwy fragment lub zbudować błędną interpretację. Jednak bez nich ryzyko konfabulacji rośnie gwałtownie. RAG nie zapewnia nieomyślności, lecz daje punkt zaczepienia.

Architektura RAG składa się z kilku etapów, które w materiałach marketingowych bywają przedstawiane jako prosta ścieżka: podziel dokumenty, stwórz embeddingi, zapisz je w bazie wektorowej, wyszukaj podobne fragmenty, dołącz je do promptu i wygeneruj odpowiedź. W praktyce każdy z tych kroków jest osobnym polem decyzji projektowych. Segmentacja, osadzanie, retrywacja i wzbogacenie promptu to podstawowe etapy procesu, z których każdy może poprawić jakość systemu albo po cichu ją zniszczyć. Źle zaprojektowany RAG jest szczególnie niebezpieczny, gdyż daje fałszywe poczucie źródłowości. Użytkownik wierzy, że system korzysta z dokumentów, podczas gdy w rzeczywistości może on operować na treściach niewłaściwych, przestarzałych, źle podzielonych lub semantycznie podobnych, lecz prawnie i faktycznie nieadekwatnych.

Pierwszym krytycznym etapem jest chunking, czyli podział dokumentów na fragmenty. Wydaje się to czynnością techniczną, niemal biurową: pokrojeniem tekstu na kawałki i zapisaniem ich w bazie. Tymczasem chunking jest jedną z najważniejszych decyzji epistemologicznych w systemie RAG. To od niego zależy, co model uzna za samodzielną jednostkę sensu.

Jeżeli fragment będzie zbyt duży, retrieval zwróci nadmiar informacji i model będzie musiał oddzielić istotne zdania od szumu. Jeśli zaś będzie zbyt mały, sens zostanie przecięty – model otrzyma zdanie bez warunku, definicję bez wyjątku lub tezę bez uzasadnienia. W obu przypadkach system może wygenerować odpowiedź brzmiącą poprawnie, lecz źle zakotwiczoną.

Trafna jest tu analogia sieci rybackiej: rozmiar fragmentu to wielkość oczek, a overlap – nakładanie się fragmentów – to gęstość splotu. Zbyt duże oczka łapią zbyt wiele szumu, zbyt małe mogą przeciąć „rybę”, czyli właściwy sens wypowiedzi. Metafora ta pokazuje, że chunking nie posiada jednego uniwersalnego rozwiązania. Inaczej dzieli się komentarz prawny, inaczej dokumentację techniczną, transkrypcję rozmowy, kod źródłowy, raport finansowy czy artykuł naukowy. Tekst nie jest masą papierniczą. Ma strukturę, hierarchię, rytm argumentacji i miejsca, których nie wolno rozcinać bez utraty sensu. W dokumentach prawnych kluczowe są definicje, wyjątki, odesłania i warunki; w technicznych – zależności proceduralne, parametry, wersje i ograniczenia środowiska;

w analizach naukowych – pojęcia, hipotezy, zakres obowiązywania twierdzeń oraz relacja między argumentem a kontrargumentem. Jeżeli system RAG traktuje wszystkie te teksty tak samo, działa jak archiwista, który tnie książki nożyczkami co tysiąc znaków, bo tak ładnie wygląda w konfiguracji. Taka neutralność techniczna jest pozorna. Każda strategia chunkingu zawiera teorię tekstu. Pytanie tylko, czy architekt jest jej świadomy.

Wyszukiwanie hybrydowe zamiast zawodnej semantyki wektorowej

Kolejnym elementem są embeddingi, czyli proces przekształcania fragmentów tekstu w wektory reprezentujące znaczenie semantyczne. W dużym uproszczeniu pozwalają one systemowi odnaleźć treści podobne znaczeniowo do zapytania użytkownika, nawet jeśli nie użyto w nich tych samych słów. To ogromna przewaga nad prostym wyszukiwaniem słów kluczowych. Gdy użytkownik pyta o „ryzyko wycieku danych osobowych”, system może wskazać fragmenty dotyczące PII, anonimizacji, prywatności, prompt injection czy kontroli okna kontekstowego.

Embeddingi tworzą mapę pokrewieństw, dzięki której język przestaje być jedynie literalnym ciągiem znaków, a staje się przestrzenią znaczeń. Należy jednak unikać bezkrytycznego zachwytu. Wektor nie „rozumie” dokumentu w ludzkim sensie; reprezentuje jedynie statystyczne relacje między fragmentami. Jest znakomity w mierzeniu podobieństwa semantycznego, lecz zawodzi przy twardych identyfikatorach, hierarchii dokumentów, statusie obowiązywania, datach, wersjach, nazwach własnych czy warunkach logicznych. Czyste wyszukiwanie wektorowe nie sprawdza się w scenariuszach biznesowych wymagających precyzji słów kluczowych, filtrów atrybutowych i relacji hierarchicznych. Wektor widzi podobieństwo.

System produkcyjny potrzebuje jednak zgodności. Różnica między podobieństwem a zgodnością jest fundamentalna: dokument semantycznie podobny nie zawsze jest dokumentem właściwym. Przepis uchylony może być niemal identyczny z obowiązującym. Stara procedura bezpieczeństwa może operować tymi samymi pojęciami co nowa. Projekt umowy może przypominać wersję podpisaną, lecz różnić się jednym usuniętym wyjątkiem, który całkowicie zmienia ryzyko prawne. Notatka techniczna z 2023 roku może brzmieć podobnie do tej z 2026 roku, mimo że dotyczy innej wersji API. Wyszukiwanie semantyczne pozbawione metadanych i filtrów odnajdzie tekst „bliski”, ale niekoniecznie „prawdziwy dla danej sprawy”. W zastosowaniach instytucjonalnych bliskość bywa zdradliwa. Bliski dokument jest jak bliski kuzyn w genealogii prawa: podobny, ale niekoniecznie dziedziczy majątek.

Stąd rekomendacja: Integrated Knowledge Retrieval, czyli wyszukiwanie hybrydowe łączące semantykę wektorową, słowa kluczowe, metadane oraz – w razie potrzeby – grafy wiedzy. Takie podejście odpowiada realnej naturze wiedzy organizacyjnej. Część zapytań wymaga podobieństwa znaczeniowego, inne dokładnego terminu, daty, statusu dokumentu, relacji między podmiotami czy hierarchii źródeł. Nie ma powodu, by zmuszać bazę wektorową do udawania całej epistemologii przedsiębiorstwa. Wektor jest świetnym narzędziem, ale nie zastąpi jednocześnie księgi wieczystej, systemu legislacyjnego, katalogu wersji i grafu zależności. Wyszukiwanie hybrydowe działa jak rozsądny bibliotekarz, który nie tylko kojarzy tematy, ale sprawdza sygnaturę, datę wydania, autora, poprawki oraz dział biblioteki. Gdy użytkownik pyta o konkretną procedurę, system powinien dostarczyć nie tylko semantycznie podobne fragmenty, lecz właściwą wersję dokumentu, akt nadrzędny, powiązane wyjątki i ramy czasowe jego obowiązywania. W obszarze przepisów, umów, norm czy danych medycznych wyszukiwanie bez metadanych przypomina prowadzenie samochodu na podstawie krajobrazu, bez znaków drogowych. Może się udać, ale trudno nazwać to systemem.

RAG jako architektura pokory i dyscypliny danych

Trzecim elementem RAG jest system of record, czyli zewnętrzny system ewidencji. Należy tu konsekwentnie odróżnić pamięć długotrwałą modelu od trwałego źródła danych, gdyż to rozróżnienie stanowi fundament odpowiedzialnej architektury. LLM nie powinien być traktowany jako magazyn prawdy. Prawda organizacyjna powinna spoczywać w kontrolowanych repozytoriach: bazach danych, systemach dokumentów, rejestrach, repozytoriach kodu, hurtowniach danych, systemach prawnych czy katalogach wiedzy. Model ma uzyskiwać do nich dostęp poprzez kontrolowane mechanizmy retrievalu, a nie polegać na mglistej pamięci rozmowy lub domniemaniu, że „pewnie coś o tym wie”. RAG jest zatem architekturą pokory. Mówi modelowi: nie odpowiadaj z majestatu własnej płynności, lecz najpierw poszukaj informacji, a jeśli ich nie znajdziesz – przyznaj to. To proste wymaganie jest w praktyce trudne, ponieważ LLM ma naturalną tendencję do wypełniania luk językiem. Dobrze zaprojektowany RAG musi więc nie tylko dostarczyć dokumenty, ale i wymusić konkretne zachowanie w przypadku ich braku. System powinien rozróżniać odpowiedź opartą na źródłach, odpowiedź częściową, taką wymagającą doprecyzowania oraz odmowę z powodu braku podstaw. Bez tego RAG staje się jedynie dekoracyjnym listkiem figowym dla halucynacji.

Czwartym elementem jest wzbogacenie promptu. Po wyszukaniu fragmentów system dołącza je do zapytania użytkownika i instruuje model, aby odpowiedział na ich podstawie. Choć brzmi to prosto, proces ten wiąże się z istotnymi decyzjami architektonicznymi: ile fragmentów dołączyć? W jakiej

kolejności? Czy udostępnić modelowi metadane, oznaczyć źródła lub wymagać cytowania? Czy zabronić wychodzenia poza kontekst, czy dopuścić uzupełnienie wiedzą ogólną?

Kwestią kluczową jest to, czy model ma sygnalizować sprzeczności między źródłami oraz czy powinien preferować dokument najnowszy, nadrzędny, zatwierdzony lub podpisany. Te pytania nie są kosmetyczne – od nich zależy, czy system rzetelnie streszczy dokumenty, czy wyprodukuje literacką zupę z fragmentów. Szczególnie ważne jest zarządzanie sprzecznością. W organizacjach dokumenty często nie tworzą harmonijnego chóru, lecz klóćcą się rodzinę. Procedury bywają aktualizowane nierówno, prezentacje marketingowe obiecują więcej niż dokumentacja techniczna, regulaminy posiadają wersje archiwalne, a w bazach wiedzy notatki robocze współlistnieją z politykami zatwierdzonymi.

RAG musi zatem nie tylko odnaleźć tekst, ale i rozpoznać jego status. W przeciwnym razie model może połączyć sprzeczne dokumenty w gładką syntezę, która sprawia wrażenie kompromisu, a w rzeczywistości jest fałszem. Nie każdą sprzeczność należy rozpuszczać w eleganckim akapicie; czasem trzeba ją wyraźnie obnażyć. Właśnie dlatego czyszczenie danych jest warunkiem działania RAG, a nie jedynie nudnym etapem przygotowawczym. „Brudne” dane są źródłem szumu semantycznego, który zatruwa bazę wektorową i uniemożliwia poprawny retrieval, niezależnie od jakości użytego modelu.

To jedno z najważniejszych zdań w całej architekturze GenAI. Organizacje często chcą zacząć od modelu, bo jest on efektowny. Powinny jednak zacząć od danych, gdyż to one są rzeczywistością, z którą model będzie musiał pracować. Jeśli ta rzeczywistość pozostanie nieuporządkowana, model stanie się tylko eleganckim rzecznikiem bałaganu.

Higiena danych i obserwowalność jako fundamenty bezpiecznego RAG

Proces czyszczenia danych to szereg żmudnych operacji: od deduplikacji i usuwania przestarzałych dokumentów, przez wersjonowanie, ekstrakcję metadanych i normalizację formatów, aż po rozpoznawanie tabel oraz oddzielanie treści głównej od nagłówków, stopek czy reklam. To także poprawa jakości OCR, rozróżnienie dokumentów źródłowych od roboczych i fundamentalna decyzja o tym, co w ogóle powinno trafić do bazy wiedzy. Każdy z tych kroków jest mniej widowiskowy niż wybór modelu, lecz każdy może przesądzić o jakości odpowiedzi. Jeśli do bazy trafi pięć wersji tej samej polityki, model może nie rozstrzygnąć, która jest właściwa. Błędne daty

sprawia, że retrieval będzie preferował starocie w świeżym kapeluszu, a źle sparsowane tabele doprowadzą do pomylenia kolumny z komentarzem.

To nie są detale; to punkty krytyczne, w których system zaczyna kłamać bez złej woli. Dlatego RAG wymaga pełnej obserwowalności. Nie wystarczy sama odpowiedź – musimy wiedzieć, jakie dokumenty pobrano, dlaczego uznano je za trafne, jaką miały punktację i które fragmenty trafiły do promptu. Musimy rozumieć, jak model je wykorzystał i czy odpowiedź rzeczywiście wynika ze źródeł. Bez tej warstwy RAG staje się czarną skrzynką z doklejonymi cytatami. Sam cytat nie jest bowiem dowodem na to, że odpowiedź została na nim oparta; model może przytoczyć właściwy dokument, wyciągając z niego błędny wniosek, lub odpowiedzieć poprawnie, cytując fragment jedynie luźno powiązany. Niezbędne są więc rygorystyczne ewaluacje: testy trafności retrievalu, groundedness, zgodności odpowiedzi ze źródłem oraz testy odmowy w przypadku braku podstaw.

W tym miejscu ujawnia się kluczowa różnica między wyszukiwarką a systemem RAG. Wyszukiwarka dostarcza listę dokumentów, którą użytkownik może samodzielnie ocenić pod kątem wiarygodności. RAG daje gotową odpowiedź – jest to rozwiązanie wygodniejsze, ale bardziej ryzykowne, gdyż redukuje możliwość kontroli. Model wykonuje część pracy za człowieka, co zwiększa produktywność, ale przesuną ciężar odpowiedzialności na architekturę. Im silniejsza synteza, tym wyraźniej system musi wskazywać źródła. W przeciwnym razie staje się maszyną do produkowania autorytatywnych skrótów, a skrót bez źródła to jedynie elegancka forma zależności od cudzej pomyłki.

RAG należy zatem projektować w oparciu o poziomy ryzyka. W zastosowaniach niskiego ryzyka – jak wstępne streszczenia czy obsługa FAQ – można dopuścić większą swobodę językową. Przy średnim ryzyku, obejmującym analizę umów, compliance czy dokumentację techniczną, wymagane są twarde źródła, wskazywanie niepewności i walidacja strukturalna. W obszarach wysokiego ryzyka – decyzjach prawnych, medycznych, finansowych czy bezpieczeństwie infrastruktury – model może być wyłącznie wsparciem, a nie samodzielnym decydentem. Architektura RAG nie znosi próżni w odpowiedzialności; powinna ją precyzyjnie umiejscowić. Na koniec pozostaje kwestia aktualności: nawet najlepszy RAG zawiedzie, jeśli baza wiedzy będzie źle utrzymywana.

ETL jako układ krwionośny RAG i wymóg higieny informacyjnej

Pipeline aktualizacji dokumentów musi być integralną częścią systemu. Nowy dokument powinien zostać pobrany, oczyszczony, podzielony, osadzony, opisany metadanymi, zapisany, przetestowany

i włączony do procesu retrievalu. Dokument wycofany musi zostać oznaczony lub usunięty, a wersje powinny być od siebie rozróżnialne. W przeciwnym razie system zacznie generować odpowiedzi z przeszłości, podając je jednak w czasie teraźniejszym. A nic tak nie dodaje fałszowi wiarygodności jak czas teraźniejszy i pewny ton.

ETL nie jest zapleczem RAG. Jest jego układem krwionośnym. To właśnie proces Extract, Transform, Load decyduje o tym, jaka wiedza w ogóle dotrze do modelu. Choć można wspomnieć o narzędziach takich jak Airbyte czy bazach wektorowych typu Pinecone, ważniejsza od nazw jest sama logika: dane muszą być pozyskiwane z kontrolowanych źródeł, transformowane z poszanowaniem struktury dokumentu i ładowane do systemu w sposób powtarzalny. Ręczne wrzucanie plików do bazy może wystarczyć do pokazu, ale nie do produkcji. Produkcja wymaga procesu, a proces odpowiedzialności.

Należy jednak pamiętać, że RAG nie jest jedyną odpowiedzią na ograniczenia LLM. Czasem właściwsze będzie klasyczne wyszukiwanie, reguły biznesowe, system ekspercki, fine-tuning, narzędzie obliczeniowe, baza grafowa lub zwykły formularz. Dojrzała architektura nie polega na wpychaniu wszystkiego do RAG, lecz na stosowaniu go tam, gdzie niezbędna jest językowa synteza nad kontrolowanym korpusem wiedzy. Jeśli system ma obliczyć podatek według znanego wzoru, model językowy nie powinien „rozumować” nad wynikiem; może go jedynie wyjaśnić po tym, jak zostanie on wyliczony przez deterministyczny moduł. Komponent od języka nie musi udawać kalkulatora, tak jak poeta nie musi obsługiwać windy.

Właściwie zaprojektowany RAG przyjmuje formę układu ról. System ewidencji przechowuje prawdę organizacyjną, ETL przygotowuje dane, a chunking dzieli tekst na jednostki sensu. Embeddings tworzą mapę podobieństw, wyszukiwanie hybrydowe znajduje właściwe fragmenty, a prompt kontekstowy przekazuje je modelowi. LLM generuje odpowiedź, walidator sprawdza zgodność formy i treści, a mechanizmy obserwowalności zapisują ścieżkę. Człowiek przejmuje decyzję tam, gdzie wymaga tego ryzyko. W takim układzie każdy element ma swoją funkcję i ograniczenia; model jest ważny, ale nie samotny. Samotność modelu to jedno z głównych źródeł halucynacji.

Największa wartość RAG ujawnia się wtedy, gdy system nie tylko odpowiada, lecz uczy organizację porządku wiedzy. Wdrożenie RAG wymaga bowiem odpowiedzi na pytania: gdzie są nasze dokumenty? Które są aktualne? Kto jest właścicielem treści? Jak oznaczamy wersje i jakie metadane mają znaczenie? Jak odróżnić dokument zatwierdzony od roboczego? Kiedy dokument wygasa, jak obsługujemy sprzeczności i jak mierzymy trafność wyszukiwania?

Wiele organizacji unikało tych pytań latami. GenAI zmusza je, by wreszcie na nie odpowiedziały. Może to być jedna z największych korzyści tej technologii: nie sama automatyzacja odpowiedzi, lecz wymuszenie higieny informacyjnej. LLM nie posiada własnego świata. RAG buduje mu okno na świat, ale jakość tego okna zależy od szkła, ramy, kierunku i tego, czy ktoś regularnie je myje. Brudne okno nie pokazuje rzeczywistości, lecz smugę. Zbyt małe okno ukazuje fragment bez kontekstu. Okno skierowane w złą stronę prezentuje właściwy krajobraz, ale z niewłaściwego budynku. A okno bez podpisu może sprawić, że użytkownik pomyli archiwum z teraźniejszością.

RAG jako praktyka epistemiczna i rola jakości danych

Architektura RAG jest zatem nie tyle funkcją techniczną, co praktyką epistemiczną. Jej celem jest zarządzanie warunkami prawdziwości w systemie, który sam z siebie nie odróżnia prawdy od prawdopodobnej wypowiedzi. Choć brzmi to abstrakcyjnie, sprowadza się do konkretnych decyzji: jakie źródła uznajemy, jak je dzielimy i opisujemy, w jaki sposób wyszukujemy, cytujemy i weryfikujemy oraz – co kluczowe – kiedy mówimy „nie wiem”. Każdy dojrzały system GenAI powinien potrafić wypowiedzieć to ostatnie zdanie. Model, który zawsze odpowiada, jest mniej użyteczny od systemu potrafiącego odmówić odpowiedzi przy braku podstaw. W świecie nadmiaru języka zdolność do milczenia staje się funkcją bezpieczeństwa.

RAG nie kończy zatem problemu, lecz stanowi początek poważnej architektury. Za nim stoją potoki ETL, decyzje o danych, wzorce integracyjne, kolejki, wersjonowanie promptów, testy, monitoring oraz polityki bezpieczeństwa i odpowiedzialność człowieka. Dopiero w takim ekosystemie model językowy może przejść z roli błyskotliwego improwizatora do roli kontrolowanego asystenta. Nie stanie się przez to świadomy – nie musi. Wystarczy, że będzie użyteczny, sprawdzalny i ograniczony. W technologii, podobnie jak w polityce, największe szkody wyrządzają nie ci, którzy mają ograniczenia, lecz ci, którzy udają, że ich nie mają.

Dane jako infrastruktura poznawcza: ETL, czyszczenie, szum semantyczny i zatrucie bazy wiedzy

RAG jest kontrolowanym oknem, przez które model językowy patrzy na świat organizacji; dane są krajobrazem za tym oknem. Można dysponować najlepszą szybą, najdroższym modelem, sprawną bazą wektorową i eleganckim promptem, ale jeśli za oknem znajduje się magazyn nieopisanych pudeł, starych procedur, duplikatów, roboczych notatek i przestarzałych regulaminów bez właściciela, system zobaczy chaos. A następnie ten chaos zostanie przedstawiony użytkownikowi językiem tak płynnym, że przez chwilę będzie sprawiał wrażenie porządku. To jedna z największych pułapek GenAI: technologia potrafi nadać bałaganowi styl.

W klasycznej informatyce jakość danych zawsze była ważna, lecz w systemach GenAI staje się kwestią zasadniczą. Model językowy nie tylko odczytuje dane – on je interpretuje, streszcza, łączy i przekształca. Jeśli otrzyma informacje błędne lub sprzeczne, niekoniecznie zatrzyma się z komunikatem ostrzegawczym. Może stworzyć syntetyczną odpowiedź formalnie poprawną i stylistycznie atrakcyjną, lecz merytorycznie fałszywą. Złe dane w systemie GenAI są zatem bardziej niebezpieczne niż w klasycznej bazie danych.

Klasyczna baza zwróci błędny rekord. LLM może zwrócić błędny rekord z uzasadnieniem, kontekstem i aksamitną pewnością głosu.

Czyszczenie danych jest etapem krytycznym przed migracją do bazy wektorowej. „Brudne” dane generują szum semantyczny, który powoduje błędne mapowanie znaczeń w przestrzeni wektorowej. W konsekwencji baza zostaje zatruta, a poprawny retrieval staje się niemożliwy niezależnie od jakości modelu.

To rozpoznanie stanowi jeden z najważniejszych mandatów architektonicznych całego tekstu. Nie istnieje model tak potężny, by unieważnić zły proces przygotowania danych. Nie ma algorytmicznego sakramentu, który rozgrzesza bałagan informacyjny. Szum semantyczny jest pojęciem szczególnie istotnym, gdyż pokazuje, że problemem nie jest zwykły „błąd”, który można znaleźć i poprawić. Szum jest bardziej podstępny; powstaje, gdy do systemu trafia zbyt wiele treści podobnych, niejednoznacznych, zduplikowanych lub pozbawionych kontekstu. Model embeddings rozmieszcza wtedy fragmenty w przestrzeni znaczeń w sposób, który formalnie działa, ale praktycznie zaciemnia relacje. Dokument roboczy łąduje blisko obowiązującego, stara procedura przypomina nową, a notatka ze spotkania miesza się z polityką organizacyjną. System nie dysponuje sygnałami pozwalającymi odróżnić wiedzę od śladu po procesie jej tworzenia.

W tym momencie dane przestają być jedynie „zasobem”, a stają się infrastrukturą poznawczą. Organizacje często mówią o danych jak o paliwie, lecz w systemach GenAI trafniejsza jest metafora gleby. Na zatrutej glebie rośnie zatruty plon, nawet jeśli użyjemy najnowocześniejszego traktora.

Dane są środowiskiem, z którego model pobiera znaczenia. Jeśli środowisko jest zdegradowane, odpowiedzi modelu będą takie same. Co gorsza, degradacja ta może być początkowo niewidoczna, gdyż język modelu przykryje ją płynnością. W klasycznym systemie problem danych ujawnia się w raporcie lub błędnym wyniku; w GenAI objawia się w narracji, która wygląda wiarygodnie. To jak pleśń pod świeżą farbą.

ETL jako fundament bezpiecznego obiegu danych w GenAI

Dlatego potoki ETL (Extract, Transform, Load) powinny być traktowane jako centralny element architektury GenAI, a nie techniczna kuchnia, do której zagląda się dopiero wtedy, gdy danie śmierdzi. Choć w praktyce ETL wiąże się z narzędziami takimi jak Airbyte czy bazami wektorowymi typu Pinecone, istotą nie są konkretne nazwy. Kluczowy jest powtarzalny, kontrolowany i audytowalny proces prowadzenia danych od źródła do formy, w której mogą być bezpiecznie wykorzystane przez system RAG. Extract to nie tylko pobranie pliku, lecz decyzja o tym, skąd wolno czerpać wiedzę. Transform to coś więcej niż zmiana formatu – to nadanie strukturze sensu. Load z kolei nie jest zwykłym zapisem do bazy, lecz wprowadzeniem fragmentu do obiegu poznawczego organizacji.

Pierwszym pytaniem procesu ETL w systemach GenAI jest pytanie o źródło. Skąd pochodzi dokument? Czy jest zatwierdzony i aktualny? Kto jest jego właścicielem i czy obowiązuje w całej organizacji, czy tylko w jednym dziale? Czy mamy do czynienia z wersją finalną, projektem, notatką, regulaminem, umową, a może jedynie luźnym zrzutem wiedzy? W klasycznych repozytoriach te kwestie bywają ignorowane, ponieważ człowiek rozpoznaje status dokumentu po kontekście. Model takiego kontekstu nie posiada, dopóki go nie otrzyma. Dla niego tekst pozostaje tekstem, dopóki metadane nie określą jego roli w instytucji.

Metadane nie są więc dodatkiem, lecz warunkiem odpowiedzialnego retrievalu. Data utworzenia i obowiązywania, autor, status zatwierdzenia, wersja, język, dział czy klasa poufności – wszystkie te informacje mogą rozstrzygać o tym, czy fragment powinien zostać użyty. W systemach prawnych lub compliance różnica między „projektem” a „wersją podpisaną” może oznaczać granicę między poradą użyteczną a szkodliwą. W systemie technicznym rozbieżność między dokumentacją wersji 1.3 a 2.0 może doprowadzić do awarii, a w HR różnica między polityką lokalną a globalną może skutkować naruszeniem prawa pracy. Semantyczna bliskość nie wystarczy; niezbędna jest instytucjonalna tożsamość dokumentu.

Drugim pytaniem jest kwestia transformacji. Dokumenty rzadko są gotowe do bezpośredniego użycia w RAG. Zawierają nagłówki, stopki, tabele, przypisy, śledzenie zmian czy błędy OCR. Jeśli ten materiał zostanie bezrefleksyjnie przekształcony w tekst, system embeddings zacznie uczyć się nie tylko znaczeń, ale i szumów. Stopka z nazwą firmy powtarzana na każdej stronie może stać się sztucznym sygnałem podobieństwa, błędnie odczytana tabela zmieni sens liczb, a komentarz roboczy zostanie uznany za obowiązujące stanowisko. Śledzone usunięcia mogą ożyć jak zombie w bazie wiedzy. A zombie dokumentacyjne są gorsze od filmowych, bo mają format PDF.

Transformacja musi zatem obejmować czyszczenie strukturalne i semantyczne. To pierwsze usuwa elementy techniczne, które nie niosą znaczenia lub zakłócają analizę. Czyszczenie semantyczne wymaga trudniejszych decyzji: czy zachować komentarze, jak oddzielić wersje archiwalne, w jaki sposób przekształcić tabele dla modelu i czy łączyć definicje z miejscami ich użycia. Musimy rozstrzygnąć, czy dzielić dokument według nagłówków i artykułów, czy jedynie według limitu tokenów.

Baza wektorowa wymaga rygorystycznego governance i walidacji

To nie jest praca „przed AI”. To praca, która decyduje o tym, czy model będzie miał z czego sensownie korzystać. Trzecim kluczowym zagadnieniem jest proces ładowania danych. Baza wektorowa nie jest śmietnikiem na teksty. Powinna być uporządkowanym rejestrem fragmentów o znanym pochodzeniu, statusie i relacji do dokumentu nadrzędnego. Każdy chunk musi być powiązany z dokumentem źródłowym, jego wersją, sekcją, datą, statusem oraz polityką dostępu. Bez tego retrieval staje się wyszukiwaniem semantycznym pozbawionym pamięci instytucjonalnej. Model może odnaleźć właściwy fragment, ale system nie będzie wiedział, czy wolno go użyć, czy jest aktualny, czy powinien być widoczny dla danego użytkownika i czy ma pierwszeństwo przed innym wynikiem. Sama baza wektorowa nie rozwiązuje problemu governance; bez niego przypomina bibliotekę pozbawioną katalogu, bibliotekarza i zakazu wynoszenia rękopisów.

W tym miejscu należy podkreślić znaczenie deduplikacji. Duplikaty to jedno z najbardziej lekceważonych źródeł szumu semantycznego. Gdy ten sam dokument występuje w kilku kopiach, system może nadać mu nieproporcjonalną wagę, ponieważ podobne fragmenty pojawiają się częściej. Jeśli różne wersje dokumentu są niemal identyczne, embeddingi umieszczają je blisko siebie, co może prowadzić do pobrania niewłaściwej wersji. Z kolei kopia robocza o treści zbliżonej do finalnej może sprawić, że model stworzy hybrydę odpowiedzi, której nigdy nie zatwierdził człowiek. W klasycznym zarządzaniu dokumentami duplikat to bałagan; w RAG jest aktywnym zakłóceniem poznawczym.

Równie istotne jest zarządzanie wersjami. W systemach GenAI wersja dokumentu musi być traktowana tak poważnie, jak wersja kodu. Skoro prompty są wersjonowane jako kod, dane zasilające RAG również powinny podlegać kontroli wersji, posiadać status i historię zmian. Logika „prompts as code” – obejmująca repozytoria, wersjonowanie semantyczne, code review i możliwość rollbacku – znajduje pełne zastosowanie w danych. O ile zmiana promptu może spowodować regresję, o tyle zmiana korpusu wiedzy bywa bardziej podstępna: model nadal odpowiada płynnie,

ale opiera się na innym materiale. Każda większa aktualizacja bazy powinna być testowana tak, jak zmianę zachowania aplikacji.

Dobrze zaprojektowany pipeline ETL nie może ograniczać się do ekstrakcji i zapisu; musi zawierać walidację jakości. Należy weryfikować, czy dokument został poprawnie odczytany, czy tabele zachowały strukturę, czy liczba chunków mieści się w normach, a metadane są kompletne. Trzeba sprawdzać brak duplikatów, aktualność dokumentu, właściwą klasyfikację poufności oraz poprawność generowania embeddingów i odzyskiwania fragmentów w zapytaniach testowych.

RAG potrzebuje własnych testów jednostkowych i integracyjnych. Jeśli organizacja nie testuje bazy wiedzy, robi to za nią użytkownik. A użytkownik w produkcji to najdroższy tester, bo wystawia fakturę w postaci utraty zaufania. Szczególną rolę odgrywają tu testy retrievalu – zestaw pytań kontrolnych z określonymi, oczekiwanymi fragmentami. Jeśli po zmianie chunkingu, modelu embeddingów lub parametrów wyszukiwania system przestaje znajdować właściwe dane, musi zostać to wykryte przed wdrożeniem. Ground truth w RAG nie jest luksusem badaczy benchmarków, lecz polisą ubezpieczeniową przed cichą degradacją systemu.

Cicha degradacja jest groźna, ponieważ użytkownik końcowy widzi jedynie odpowiedź, a nie proces retrievalu. Jeśli brzmi ona rozsądnie, może jej zaufać, mimo że opiera się na drugim najlepszym fragmencie (bo pierwszy nie został znaleziony), dokumencie nieaktualnym lub cytacie wyrwanym z tabeli bez kontekstu kolumn. Model nie „wie”, że kontekst jest niepełny – on go po prostu wykorzysta.

Dlatego system powinien mierzyć nie tylko jakość odpowiedzi, ale i drogę, która do niej prowadzi. W dojrzałym RAG każde pytanie powinno zostawiać ślad: od pierwotnego zapytania, przez jego transformację i znalezione fragmenty z punktacją, aż po zastosowane filtry, finalny prompt i reakcję użytkownika. Taka obserwowalność nie jest biurokracją, lecz warunkiem uczenia się systemu. Bez niej organizacja nie wie, czy problem leży w modelu, danych, chunkingu czy metadanych, przez co zaczyna „tuningować” na oślep. Tuning na oślep to rytuał deszczu w chmurze obliczeniowej.

Warto również rozróżnić dwa rodzaje jakości: formalną i epistemiczną. Jakość formalna oznacza, że dokument jest poprawnie odczytany, zakodowany i opisany. Jakość epistemiczna to wiarygodność, aktualność i adekwatność treści. Można mieć dane formalnie perfekcyjne, ale epistemicznie bezużyteczne – stary regulamin może być pięknie sparsowany, a błędna notatka może mieć idealne embeddingi. System GenAI potrzebuje obu tych jakości jednocześnie: bez formalnej nie znajdzie sensu, a bez epistemicznej znajdzie sens nieprawdziwy.

Z tego powodu projekty GenAI wymagają współpracy inżynierów danych, ekspertów domenowych, prawników, specjalistów bezpieczeństwa i UX. To przedsięwzięcie multidyscyplinarne, łączące kompetencje prompt engineerów, data engineers oraz badaczy QA i UX.

Interdyscyplinarny zespół jako gwarant bezpieczeństwa wdrożenia

To nie jest ozdobny skład zespołu. To konieczność. Inżynier danych może wiedzieć, jak oczyścić dokument. Ekspert domenowy wie, czy dokument powinien być używany. Prawnik wie, czy można przetwarzać dane osobowe. Specjalista bezpieczeństwa wie, co grozi wyciekiem. UX researcher wie, jak użytkownik zinterpretuje odpowiedź. QA wie, jak wykryć regresję. Sam model nie zastąpi żadnej z tych odpowiedzialności. Może je najwyżej zamaskować.

Interdyscyplinarny pipeline danych jako fundament bezpieczeństwa

W praktyce organizacje często próbują pominąć tę interdyscyplinarność, uznając ją za zbyt kosztowną. Tymczasem jej brak kosztuje jeszcze więcej. Projekt GenAI realizowany wyłącznie przez entuzjastów modelu może szybko przynieść efekt, ale równie szybko wpadnie w bagno danych, uprawnień, wersji, wyjątków, kosztów i odpowiedzialności. Z kolei inicjatywa prowadzona jedynie przez tradycyjny dział IT może być bezpieczna, lecz nie wykorzysta w pełni potencjału językowego modeli. Projekt bez udziału ekspertów domenowych z kolei może działać poprawnie technicznie, odpowiadając jednocześnie błędnie merytorycznie.

GenAI jest technologią graniczną: operuje na styku języka, danych, procesów i instytucji. Wymaga zatem zespołu rozumiejącego więcej niż jedną warstwę rzeczywistości. Szczególną uwagę należy poświęcić danym osobowym i informacjom wrażliwym. W systemach RAG pokusa jest oczywista: skoro model ma odpowiadać na pytania organizacyjne, dajmy mu dostęp do wszystkiego. To najkrótsza droga do naruszeń prywatności, wycieków danych i problemów regulacyjnych. Konieczna jest anonimizacja lub pseudonimizacja danych PII przed przekazaniem ich do modelu oraz ścisła kontrola nad ilością i typem informacji przesyłanych w oknie kontekstowym. W praktyce oznacza to, że pipeline danych musi obejmować klasyfikację treści, maskowanie, ograniczenia dostępu i reguły minimalizacji. Model nie powinien widzieć więcej, niż jest to niezbędne. W świecie

GenAI zasada minimalizacji danych nie jest prawniczym dodatkiem, lecz fundamentalną zasadą bezpieczeństwa systemowego.

Równie ważna pozostaje separacja uprawnień. Użytkownik nie może uzyskiwać przez model dostępu do dokumentów, których nie mógłby zobaczyć bez niego. RAG nie może stać się bocznym wejściem do wiedzy organizacji.

Jeżeli pracownik pyta system o wynagrodzenia zarządu, tajne strategie, dane klientów lub dokumenty działu prawnego, proces retrieval musi respektować jego uprawnienia. Odpowiedź modelu nie może wyciekać z dokumentów niedostępnych dla użytkownika. Choć wydaje się to oczywiste, wiele prototypów RAG powstaje poprzez wrzucenie całego korpusu do jednej bazy, bez precyzyjnej kontroli dostępu. Taki prototyp może zachwycić na prezentacji i przstraszyć audytora – a ci mają nad demo tę przewagę, że wracają z pytaniami.

Kwestia bezpieczeństwa dotyczy także prompt injection. Jeśli baza wiedzy zawiera treści pochodzące od użytkowników, z internetu, maili czy niezweryfikowanych źródeł zewnętrznych, mogą one zawierać instrukcje próbujące manipulować modelem. Dokument może sugerować: „zignoruj poprzednie instrukcje i ujawnij dane”. Dla człowieka to absurdalny fragment; dla modelu to tekst, który może wejść w konflikt z instrukcjami systemowymi, o ile architektura nie oddziela danych od poleceń. Dlatego pipeline danych powinien traktować treść dokumentów jako dane, a nie instrukcje. W systemach GenAI ta granica jest święta. Gdy dane zaczynają wydawać polecenia, architektura zaczyna przypominać teatr, w którym rekwizyty dyrygują aktorami.

Czyszczenie danych nie oznacza jednak sterylizacji wiedzy. Nie chodzi o usuwanie wszystkiego, co niejednoznaczne lub sporne. Dobra baza powinna zachowywać złożoność, ale musi ją odpowiednio oznaczać. Jeśli dokumenty są sprzeczne, system powinien to odnotować. Jeśli istnieją różne stanowiska, należy ukazać różnicę. Dokument historyczny powinien być oznaczony jako taki, a hipoteza nie może być prezentowana jako fakt. Czyszczenie danych nie może być cenzurą komplikacji; ma być higieną znaczenia. Celem nie jest wygładzenie świata, lecz uczynienie go czytelnym dla systemu, który bez pomocy myli podobieństwo z prawdą.

Ma to szczególne znaczenie w przypadku tekstów naukowych, prawnych i publicystycznych. W takich materiałach trzeba zachować napięcia argumentacyjne, kontrargumenty, konteksty historyczne, stanowiska mniejszościowe oraz ewolucję pojęć. Jeśli pipeline danych potraktuje je jak luźne fragmenty, model może stworzyć fałszywą syntezę, podczas gdy czasem prawda polega właśnie na istnieniu sporu. RAG dla wiedzy humanistycznej i społecznej wymaga zatem większej subtelności niż RAG dla prostego FAQ. Nie wystarczy znaleźć podobny akapit – trzeba rozumieć funkcję, jaką pełni on w argumentacji: czy jest to definicja, przykład, krytyka, ironia, czy może

wniosek? Bez tej warstwy model może cytować polemikę jako stanowisko autora, a wyjątek jako regułę. To dowodzi, że ETL dla GenAI jest także pracą hermeneutyczną.

GenAI jako lustro organizacji i wymóg dyscypliny danych

Oczywiście nie mówię tu o romantycznej kontemplacji tekstu przy świecy, choć ta czasem pomaga bardziej niż sprint bez wymagań. Chodzi o to, że transformacja danych musi uwzględniać znaczenie, strukturę i funkcję poszczególnych fragmentów. W prostych zastosowaniach wystarczy podział na akapity. W złożonych trzeba rozpoznawać sekcje, definicje, relacje, argumenty, zależności, tabele, metody, wyniki, zastrzeżenia, status dokumentu oraz jego miejsce w większym korpusie. Im bardziej odpowiedzialne zadanie, tym mniej wolno traktować tekst jak anonimową masę tokenów.

Projektowanie bazy wiedzy dla GenAI przypomina budowę instytucji. Potrzebne są reguły przyjmowania dokumentów, ich klasyfikacji, aktualizacji i wycofywania, a także zasady dostępu, audytu, rozstrzygania konfliktów oraz odpowiedzialności. Bez nich system RAG będzie działał, ale jak instytucja pozbawiona kancelarii, archiwum i procedury obiegu dokumentów.

Coś znajdzie, coś odpowie, coś wygeneruje. Pytanie brzmi: czy będzie można temu zaufać, gdy stawką nie jest ciekawostka, lecz decyzja? Dane są więc pierwszym miejscem, w którym obietnica GenAI spotyka się z kulturą organizacyjną. Organizacja uporządkowana może użyć LLM jako akceleratora wiedzy. Organizacja chaotyczna użyje go jako akceleratora chaosu. To brutalna diagnoza, ale uczciwa. Model nie naprawi dokumentów, których nikt nie chce uporządkować. Nie rozstrzygnie, która procedura obowiązuje, jeśli organizacja sama tego nie wie. Nie usunie sprzeczności między działami, jeśli jest ona polityczna, a nie techniczna. GenAI może być lustrem, ale lustro nie sprząta pokoju. Co najwyżej pokaże, że kurz nauczył się mówić.

Dlatego odpowiedzialne wdrożenie powinno zaczynać się od audytu danych. Jakie korpusy mają zasilać system? Jakie są ich źródła i jak często się zmieniają? Kto je zatwierdza? Jakie zawierają dane wrażliwe, jakie mają formaty i jakie są znane problemy jakościowe? Jak mierzymy aktualność? Jak rozpoznajemy dokument nadrzędny? Jakie metadane są obowiązkowe? Jakie pytania użytkowników system ma obsługiwać, a jakich nie?

Dopiero po odpowiedziach na te pytania warto wybierać bazę wektorową, model embeddings i parametry chunkingu. Odwrotna kolejność jest częsta, bo narzędzia są kuszące, ale architektura zaczynająca się od narzędzia zwykle kończy się obejściem problemu, a nie jego rozwiązaniem. Nie

oznacza to jednak, że każdy projekt musi startować od wielomiesięcznej reformy danych. Notatka słusznie zaleca krótkie POC, nawet dwudniowe, by szybko weryfikować założenia i nie marnować tygodni na ślepe zaułki. Krótkie POC nie powinno jednak udawać produkcji; jego celem jest wykrycie, czy dane, proces i model mają sensowny punkt styku. Jeśli POC pokazuje, że retrieval działa tylko na trzech starannie wybranych dokumentach – to jest cenna informacja. Jeśli wykazuje, że dokumenty są nieczytelne, sprzeczne albo zbyt brudne, to również jest sukces poznawczy. POC obnażający niewykonalność oszczędza pieniądze. W świecie hype'u porażka po dwóch dniach jest często zwycięstwem nad porażką po sześciu miesiącach.

Najdojrzałe organizacje potraktują GenAI jako impuls do przebudowy zarządzania wiedzą. Może to być paradoksalnie ważniejsze niż sama automatyzacja odpowiedzi. Aby model odpowiadał dobrze, organizacja musi wiedzieć, co wie. Musi odróżnić dokument od plotki, procedurę od praktyki, wersję od kopii, źródło od komentarza, fakt od hipotezy oraz archiwum od teraźniejszości. To klasyczne zadania instytucji poznawczej, których pilność GenAI tylko zwiększa. Można powiedzieć, że sztuczna inteligencja wymusza naturalną inteligencję organizacyjną.

I bardzo dobrze, bo w wielu firmach i urzędach była ona dotąd trzymana w folderze „do uporządkowania kiedyś”. W ten sposób kwestia danych domyka wcześniejszą część o RAG. Retrieval-Augmented Generation nie zaczyna się w chwili, gdy użytkownik zadaje pytanie. Zaczyna się znacznie wcześniej: przy decyzji, jakie dokumenty stanowią prawdę organizacji, kto je utrzymuje, jak są czyszczone, dzielone, osadzane, wersjonowane i wycofywane. Model językowy jest widoczną twarzą systemu, ale twarz ta mówi tym, czym została nakarmiona. Jeżeli karmimy ją śmieciami, nie dziwny się, że śmieci zaczną przemawiać stylem eksperta.

Pattern-Guided Coding jako przejście od vibe codingu do architektury

Dopiero na tak przygotowanym fundamencie można przejść do Pattern-Guided Coding – metody, która porządkuje nie tylko dane, ale i sam proces tworzenia systemów GenAI. Skoro model jest zawodnym komponentem, RAG jest protezą źródłowości, a dane stanowią infrastrukturę poznawczą, niezbędny staje się język projektowania pozwalający budować systemy powtarzalne, audytowalne i odporne. Tym językiem są wzorce: GoF, Enterprise Integration Patterns, adaptery, strategie, komendy, kolejki, bezpieczniki czy kanały błędów. Trzeba przejść od „vibe codingu” do architektury, w której każde zakłęcie ma wzorzec, a każdy wzorzec ma swoją odpowiedzialność.

Pytanie o to, jak właściwie budować systemy GenAI, nie dotyczy tego, jak je prezentować na slajdach czy sprawić, by przez pięć minut wyglądały jak przyszłość pracy. Chodzi o to, jak tworzyć rozwiązania powtarzalne, zrozumiałe, skalowalne i odporne na awarie. W tym miejscu pojawia się Pattern-Guided Coding – kodowanie oparte na wzorcach, będące metodą redukcji chaosu nieustrukturyzowanego promptowania poprzez użycie precyzyjnego języka wzorców projektowych.

To przesunięcie jest kluczowe. W pierwszej fazie fascynacji GenAI wielu odkryło, że można „rozmawiać” z modelem i otrzymywać kod, diagramy czy konfiguracje. Powstała kultura „vibe codingu”: człowiek opisuje intencję, model generuje rozwiązanie, a następnie obie strony iteracyjnie dopracowują efekt. W zastosowaniach eksperymentalnych taka praktyka jest twórcza – pozwala prototypować i przełamywać próg wejścia. Problem zaczyna się jednak wtedy, gdy improwizacja zostaje pomyślona z inżynierią. Vibe coding dobrze znosi demo. Produkcja ma słabsze poczucie humoru.

Pattern-Guided Coding nie neguje użyteczności modeli w generowaniu kodu; przeciwnie, próbuje wydobyć z niej maksimum, ograniczając ryzyko. Podstawowa intuicja jest prosta: skoro LLM generuje lepsze wyniki, gdy otrzymuje precyzyjne struktury i język znany z dobrze udokumentowanych korpusów, należy komunikować się z nim właśnie językiem wzorców. Wzorce projektowe są dla inżynierii oprogramowania tym, czym pojęcia konstytucyjne dla prawa lub kategorie metodologiczne dla nauki: nie zastępują myślenia, ale umożliwiają myślenie zbiorowe. Pozwalają nazwać problem, porównać rozwiązania i przenieść doświadczenie między różnymi sytuacjami.

PGC wykorzystuje klasyczne wzorce inżynierskie jako fundament projektowania systemów GenAI. Zamiast traktować „wzorce agentowe” jako zupełnie nowy byt, metoda mapuje je na sprawdzone rozwiązania GoF i EIP. Agent przestaje być tajemniczą figurą autonomii, a staje się kompozycją komponentów: adapterów, strategii, komend, kolejek, mechanizmów retry, bezpieczników i kanałów błędów. Ta translacja ma fundamentalne znaczenie praktyczne.

Wzorce projektowe jako fundament stabilności systemów GenAI

Kiedy mówimy „agent”, nie wiadomo jeszcze, jak system obsłuży awarię. Kiedy jednak mówimy „Command z retry i Dead Letter Channel”, zaczynamy widzieć architekturę. Wzorce są pamięcią branży; każdy z nich to skondensowana odpowiedź na problem, który wielokrotnie powtarzał się w różnych systemach. Adapter powstał, bo komponenty nie mówiły tym samym językiem. Strategia

wprowadzono, by algorytm można było wymieniać bez niszczenia reszty systemu. Command pozwolił opakować żądanie w obiekt, co umożliwiło jego kolejkowanie, logowanie i odtwarzanie. Circuit Breaker zapobiega bez końca uderzaniu w usługę, która przestała odpowiadać, a Dead Letter Channel gwarantuje, że wiadomości niemożliwych do przetworzenia nie zostaną zgubione ani mielone w nieskończoność. Te wzorce są historią ran, które zamieniono w architekturę. GenAI, z jego niedeterminizmem i skłonnością do halucynacji, bardzo potrzebuje historii cudzych ran.

Jedną z największych zalet PGC jest przełamywanie barier komunikacyjnych. W zespołach wdrażających GenAI spotykają się specjaliści posługujący się różnymi językami: data scientists, inżynierowie backendu, architekci, eksperci domenowi, prawnicy, UX researcherzy, product ownerzy i menedżerowie. Bez wspólnego słownika rozmowa szybko zamienia się w targ metafor.

Jedni mówią o agentach, inni o endpointach, a jeszcze inni o workflow, automatyzacji, pamięci, narzędziach, promptach czy doświadczeniu użytkownika. PGC proponuje sprowadzenie tej dyskusji do języka wzorców: tutaj mamy komponent decyzyjny, tutaj adapter narzędzia, tutaj stan sesji, system ewidencji, brokera komunikatów, krytyka, walidatora i kanał błędów. To nie zabija kreatywności. To daje jej adres zameldowania.

Drugą zaletą jest ujawnianie decyzji projektowych. W swobodnym kodowaniu wspomagany przez LLM wiele z nich zostaje ukrytych w wygenerowanym kodzie. Model tworzy strukturę, wybiera sposób obsługi błędów i przyjmuje założenia dotyczące stanu, wejścia, wyjścia oraz kolejności działań. Jeśli zespół nie nazwie tych decyzji, trudno je ocenić. PGC wymusza ich wcześniejsze wydobywanie. Nie wystarczy powiedzieć: „zrób agenta do obsługi zapytań klientów”. Trzeba zapytać: czy przepływ ma być orkiestracją, czy choreografią? Jakie adaptory wywołuje? Jakie dane znajdują się w stanie sesji? Co stanowi system ewidencji? Które błędy są przejściowe, które logiczne, a które krytyczne? Gdzie następuje walidacja? Kiedy człowiek przejmuje decyzję? Jak wygląda rollback i jakie są limity kosztów? Na tym polega różnica między życzeniem a projektem.

Trzecia zaleta dotyczy warstwy integracji. Wiele problemów produkcyjnych w systemach GenAI nie wynika bezpośrednio ze słabości modelu, lecz z błędów infrastrukturalnych: zgubionych wiadomości, timeoutów, braku idempotencji, niekontrolowanych retry, przeciążenia API, braku izolacji awarii czy niewłaściwej obsługi komunikatów błędnych. Modele są widowiskowym elementem systemu, ale awarie najczęściej zdarzają się w hydraulice. A hydraulika jest dziedziną niewdzięczną: gdy działa, nikt jej nie chwali; gdy zawodzi, wszyscy natychmiast to zauważają. Enterprise Integration Patterns są dla GenAI szczególnie ważne, ponieważ systemy agentowe są z natury systemami integracyjnymi. Model sam w sobie niczego nie wykonuje w świecie – musi wywołać narzędzie, pobrać dane, zapisać wynik, wysłać komunikat, wykonać transformację, poprosić inny komponent o ocenę lub przekazać zadanie dalej.

Wzorce integracyjne jako zabezpieczenie przed ryzykiem

Każde takie przejście jest punktem ryzyka. Adapter może zwrócić błąd, API przekroczyć limit, a kolejka się zapchać. Worker może ulec awarii w połowie zadania, model wygenerować błędny format, a walidator odrzucić wynik. Do tego dochodzi ryzyko złośliwych danych podanych przez użytkownika. Wzorce integracyjne nie są zatem eleganckim dodatkiem akademickim.

Są mapą pola minowego. Szczególnie istotny okazuje się wzorzec Adapter. W kontekście GenAI często mówi się o „tools”, czyli narzędziach dostępnych dla modelu. Słowo to brzmi niewinnie: model posiada narzędzia, więc może działać. Jednak z perspektywy inżynierskiej każde takie narzędzie jest adapterem do zewnętrznej funkcji, API, bazy danych, usługi lub procesu. Adapter musi posiadać kontrakt wejścia i wyjścia, walidację parametrów, obsługę błędów, politykę uprawnień oraz mechanizmy ograniczania skutków ubocznych. Model nie powinien „po prostu” wywoływać narzędzia; powinien wygenerować intencję zgodną ze schematem, a adapter ma zdecydować, czy tę intencję można bezpiecznie zrealizować. To rozróżnienie chroni system przed sytuacją, w której tekstowa sugestia modelu staje się niekontrolowaną operacją w świecie rzeczywistym.

Wzorzec Command również nabiera w GenAI szczególnej wagi. Jeśli działanie agenta zostanie opakowane jako komenda, można je zapisać, zweryfikować, kolejkować, powtórzyć, odrzucić, podpisać, audytować i przekazać do wykonania innemu komponentowi. Komenda oddziela decyzję od jej realizacji. Jest to kluczowe, ponieważ decyzja modelu jest probabilistyczna, podczas gdy wykonanie operacji w systemie może mieć trwałe skutki. Gdy model sugeruje „wyślij mail”, „utwórz rekord”, „zmień status”, „zaktualizuj bazę”, „złóż płatność” lub „usuń dokument”, system nie powinien traktować tego jako natychmiastowego rozkazu. Powinien najpierw utworzyć komendę, sprawdzić ją, ewentualnie wymagać zatwierdzenia i dopiero wtedy przystąpić do wykonania.

W systemach wysokiego ryzyka ten dystans między generacją a skutkiem jest moralnym odpowiednikiem hamulca. Z kolei wzorzec Strategy pomaga tam, gdzie system musi dobierać różne ścieżki działania. Zapytanie użytkownika może być kierowane na przykład do RAG, klasycznego wyszukiwania, narzędzia obliczeniowego, bazy SQL, człowieka, modułu tłumaczenia lub zakończyć się odmową odpowiedzi. Bez Strategy logika wyboru zostaje zaszyta w promptach i rozproszona po kodzie; dzięki niej staje się jawna i wymienna.

Pozwala to testować, kiedy system wybiera model drogi, a kiedy tani, kiedy retrieval, walidator, a kiedy eskalację. Ma to również znaczenie kosztowe. System, który każde pytanie wysyła do

największego modelu z pełnym kontekstem, nie jest inteligentny. Jest rozrzutny. Czasem najlepsza strategia polega na tym, by w ogóle nie używać LLM, gdy wystarczy prosta reguła.

Wzorce projektowe i Topologos jako ramy bezpieczeństwa systemu

Wzorzec Template Method pozwala uporządkować pracę konsumentów, czyli komponentów przetwarzających zadania. W potokach GenAI wiele operacji opiera się na powtarzalnej strukturze: odebranie komunikatu, walidacja wejścia, pobranie kontekstu, wywołanie modelu lub narzędzia, weryfikacja wyniku, zapis obserwacji oraz potwierdzenie przetworzenia lub skierowanie błędu do odpowiedniej kolejki. Template Method umożliwia zachowanie jednolitego szkieletu procesu mimo różnic w szczegółach implementacyjnych. Jest to kluczowe, ponieważ systemy agentowe mają tendencję do rozrastania się w małe królestwa wyjątków. Z czasem każdy worker zaczyna działać nieco inaczej – odmiennie obsługuje błędy, loguje dane czy waliduje wejścia. Po kilku miesiącach nikt już nie wie, gdzie faktycznie mieszka logika.

Template Method jest lekiem na feudalizm mikroserwisów. Z kolei wzorzec Proxy odpowiada za bezpieczeństwo. Model nie powinien mieć bezpośredniego dostępu do wrażliwych systemów; Proxy może filtrować dane, maskować je, sprawdzać uprawnienia, ograniczać zakres informacji, rejestrować operacje i egzekwować polityki. W systemach operujących na danych osobowych, tajemnicach przedsiębiorstwa, dokumentach prawnych czy finansowych nie jest to opcja, lecz obowiązek. Prompt injection, data leakage i nadmierne uprawnienia modelu to nie egzotyczne zagrożenia z laboratoriów, ale realne ryzyka systemów, które zamieniają język w interfejs do działania. Gdy język staje się interfejsem, każdy tekst może próbować być poleceniem. Proxy przypomina, że nie każdy tekst zasługuje na posłuch.

Pattern-Guided Coding (PGC) jest zatem metodą ograniczania halucynacji nie tylko w treści odpowiedzi, ale i w samej architekturze. Model generujący kod potrafi stworzyć rozwiązanie, które wygląda sensownie, lecz zawiera ukryte braki: brak mechanizmów retry, limitów, walidacji, idempotencji, separacji danych od instrukcji, kolejki błędów czy obserwowalności. Jeśli prompt zawiera wyłącznie opis celu, model wypełni luki domyślnymi wzorcami z danych treningowych, które niekoniecznie pasują do konkretnego systemu. Jeśli natomiast prompt definiuje precyzyjne wzorce, topologię, ograniczenia i kryteria akceptacji, przestrzeń błędu maleje. Model nie staje się przez to nieomylny, ale zostaje zamknięty w lepiej opisanej klatce.

A w przypadku komponentu, który potrafi pewnie fantazjować, dobra klatka jest aktem życzliwości wobec produkcji. Ważną rolę pełni tu Topologos – rozumiany jako specyficzna umiejętność lub wielorazowy prompt, który pozwala budować bezpieczne ramy systemowe za pomocą jednej komendy i prowadzić projektowanie poprzez kolejne decyzje. Nie należy traktować go jako magicznego generatora architektury, lecz jako przykład proceduralizacji PGC. Topologos wymusza, aby przed powstaniem kodu wyjaśnić zakres problemu, skalę, niezawodność, infrastrukturę, bezpieczeństwo i tolerancję na awarie. Następnie analizuje problem przez pryzmat wzorców, przedstawia decyzje do akceptacji, wskazuje trade-offs i dopiero na końcu generuje artefakty, takie jak manifesty RabbitMQ czy konfiguracje Terraform.

Kluczowa jest tutaj iteracyjna akceptacja. W kulturze promptowania istnieje pokusa, by poprosić model o „pełne rozwiązanie” i otrzymać jeden wielki blok kodu lub konfiguracji. Jest to szybkie, ale ryzykowne; taki blok jest trudny do oceny, zwłaszcza gdy zawiera wiele ukrytych założeń. Topologos prowadzi proces krok po kroku: proponuje, uzasadnia, wskazuje wzorce oraz wady i zalety, a człowiek zatwierdza lub odrzuca propozycje. To właściwy model współpracy z AI w zadaniach architektonicznych. Model przyspiesza analizę i generowanie wariantów, ale człowiek pozostaje właścicielem decyzji. Sztuczna inteligencja może być świetnym sekretarzem projektowym, lecz nie powinna być niewybieralnym zarządem architektury.

PGC jako narzędzie audytowalności i dyscypliny wdrożeniowej

PGC ma również istotne znaczenie dokumentacyjne. Jeśli system projektowany jest w oparciu o wzorce, dokumentacja może powstawać równolegle z projektem: diagramy Mermaid lub Graphviz, manifesty infrastruktury, opisy kanałów komunikacji, topologie kolejek, reguły retry, klasy błędów, granice komponentów czy decyzje dotyczące bezpieczeństwa. PGC pozwala dokumentować architekturę za pomocą standardów, a następnie generować na ich podstawie manifesty i kod, co redukuje ryzyko halucynacji poprzez osadzenie modelu w sztywnych ramach inżynierskich. Jest to kluczowe, gdyż jednym z problemów AI-assisted development jest tempo produkcji kodu, któremu nie towarzyszy adekwatne tempo rozumienia. Kod rośnie szybciej niż dokumentacja, a dokumentacja szybciej niż odpowiedzialność. PGC próbuje odwrócić tę kolejność: najpierw decyzja, potem artefakt.

Nie należy jednak idealizować wzorców, gdyż mogą być one nadużywane. Można nimi ozdabiać banalne rozwiązania, komplikować proste systemy i budować architektoniczne katedry tam, gdzie wystarczyłby garaż. Pattern-Guided Coding nie powinien oznaczać rytualnego wpychania GoF i EIP

do każdego promptu; jego celem jest adekwatność, a nie ornament. Czasem najlepszym wzorcem jest brak wzorca, bo problem okazuje się prosty. Czasem lepszą decyzją będzie użycie klasycznego formularza zamiast agenta, a czasem RabbitMQ okaże się zbędny, jeśli zadanie jest synchroniczne, krótkie i niskiego ryzyka.

Dojrzałość inżynierska polega także na tym, by nie budować lotniska dla papierowego samolotu. Tu pojawia się istotny kontrargument wobec PGC: czy formalizacja nie spowalnia innowacji? Czy szybkie prototypowanie nie wymaga swobody i czy wzorce nie zamienią GenAI w kolejną warstwę korporacyjnej ciężkości? Odpowiedź wymaga rozróżnienia. Na etapie eksploracji swoboda jest cenna – krótkie POC, eksperymenty, szybkie szkice i modele robocze mają głęboki sens. Notatka sama zaleca tworzenie małych POC z twardym limitem czasu, aby szybko obnażać błędne założenia. Jednak to, co jest dopuszczalne w laboratorium, nie może trafić do produkcji. PGC nie musi dławić fazy odkrywania; ma chronić fazę wdrożenia. Problemem nie jest sam eksperyment, lecz eksperyment przebrany za system krytyczny.

PGC powinno być zatem skalowane do ryzyka. W projekcie niskiego ryzyka wystarczą lżejsze ramy: wersjonowany prompt, podstawowy walidator, proste testy i ograniczony dostęp do danych. Projekt średniego ryzyka wymaga już formalnych wzorców integracji, monitorowania, kontroli kosztów, ewaluacji oraz polityk bezpieczeństwa. W przypadku projektów wysokiego ryzyka niezbędna jest pełna architektura: kolejki, manual ACK, wielopoziomowe DLQ, human-in-the-loop, testy regresyjne, audyt, ścisła kontrola dostępu i formalna walidacja wyników.

Wzorce nie są ideologią, lecz skrzynką narzędziową. Profesjonalizm polega na wyborze właściwego narzędzia, a nie na noszeniu całej skrzynki na plecach przy każdej naprawie kranu. PGC jest także odpowiedzią na problem odpowiedzialności w zespołach. Jeśli architektura powstaje jako seria promptów i odpowiedzi modelu, trudno później ustalić przyczynę konkretnych decyzji projektowych. Gdy każda decyzja jest powiązana z nazwanym wzorcem, trade-offem i akceptacją człowieka, system staje się audytowalny. Można wtedy wrócić do momentu, w którym wybrano orkiestrację zamiast choreografii, retry zamiast natychmiastowej eskalacji, konkretną strategię chunkingu, model czy próg walidacji. Ma to znaczenie nie tylko techniczne, ale i prawne oraz organizacyjne. W świecie regulacji i odpowiedzialności biznesowej argument „model tak wygenerował” nie jest akceptowalnym uzasadnieniem – to cyfrowa wersja "pies zjadł pracę domową". Istotna jest również edukacyjna funkcja PGC.

Wzorce projektowe jako droga od efektywnych prototypów do stabilnych systemów

Modele językowe obniżają próg wejścia do tworzenia oprogramowania, ale mogą jednocześnie obniżać próg rozumienia, jeśli użytkownik zaczyna akceptować wygenerowane rozwiązania bez pojęcia, jak one właściwie działają. Kodowanie oparte na wzorcach wymusza naukę architektury. Użytkownik nie prosi już tylko o to, by „zrobić aplikację”, lecz uczy się pytać: jaki wzorzec integracyjny jest właściwy? Jak obsłużyć awarie? Jakie będą skutki uboczne? Jak modelować stan, separować instrukcje od danych czy ograniczać uprawnienia?

W ten sposób AI nie zastępuje kompetencji, lecz może je podnosić — pod warunkiem oczywiście, że użytkownik chce się uczyć, a nie tylko odbierać cyfrowe dania z taśmy. W szerszej perspektywie PGC można uznać za próbę cywilizowania GenAI. Każda nowa technologia przechodzi fazę dzikiego zachwyty, potem rozczarowania, a następnie — jeśli przetrwa — fazę instytucjonalizacji. Wzorce są narzędziem tej trzeciej etapy. Sprawiają, że technologia przestaje być serią cudów, a staje się praktyką. Praktyka zaś wymaga słownika, standardów, procedur, kryteriów jakości, odpowiedzialności i pamięci błędów. Bez tego GenAI pozostanie kulturą efektywnych prototypów i kosztownych wdrożeń, które gasną w momencie, gdy trzeba je utrzymać. Z wzorcami ma szansę stać się normalnym elementem systemów IT. A normalność, choć mniej seksowna od hype'u, ma tę przewagę, że działa we wtorek po południu.

Pattern-Guided Coding nie polega na zastąpieniu człowieka wzorcem, lecz na tym, by człowiek i model spotkali się w języku bardziej precyzyjnym niż intuicja, a jednocześnie bardziej elastycznym niż sztywna specyfikacja. Model może proponować warianty, generować diagramy, pisać manifesty, przygotowywać kod, wskazywać ryzyka czy porównywać trade-offs. Człowiek musi natomiast rozumieć dziedzinę, zatwierdzać decyzje, ponosić odpowiedzialność i wiedzieć, kiedy model przesadza. Wzorce są tu wspólnym stołem roboczym. Bez niego współpraca z AI przypomina układanie silnika na dywanie w salonie: przez chwilę ekscytujące, potem zostaje plama.

Pattern-Guided Coding prowadzi nas naturalnie do infrastruktury komunikatów, szczególnie RabbitMQ, oraz do implementacji wzorców agentowych, takich jak ReAct. Jeżeli agent jest mikroarchitekturą, a jego działanie opiera się na pętli myśl–akcja–obserwacja, trzeba zapytać, jak tę pętlę osadzić w systemie odpornym na awarie. Jak nie zgubić wiadomości, gdy model długo „myśli”? Jak skalować pracowników narzędziowych? Jak odróżnić błąd sieciowy od błędu logicznego modelu?

Jak ponawiać próby bez wywoływania burzy retry? Jak izolować komunikaty, których nie da się przetworzyć? Jak sprawić, by agent nie był zabawką w notebooku, lecz elementem produkcyjnego systemu rozproszonego? To właśnie będzie przedmiotem kolejnej części: ReAct, RabbitMQ, manual ACK, competing consumers i wielopoziomowe kolejki błędów jako inżynierska odpowiedź na romantyczny mit autonomicznego agenta.

Sztuczna inteligencja w organizacji działa jak lustro: może stać się potężnym akceleratorem wiedzy, ale równie dobrze może być jedynie aksamitnym głosem, który z niezwykłą pewnością siebie opakuje istniejący chaos w piękne zdania. Ostatecznie to nie moc modelu, lecz higiena danych i dyscyplina architekta decydują o tym, czy budujemy użyteczne narzędzie, czy jedynie kosztowny pomnik cyfrowej iluzji. Czy jesteśmy zatem gotowi przestać rozmawiać z AI jak z wyrocznią, a zacząć projektować ją jak system rozproszony?

Inspiracje

[1]



"Building Complex Multi-Agent Systems Using" Tim O'Brien

2026 | Packt Publishing | ISBN: 9781806114290

Tim O'Brien jest liderem technologicznym, architektem i programistą z ponad 20-letnim doświadczeniem w branży oprogramowania. Pełnił funkcję dyrektora ds. technologii (CTO) i założyciela wielu firm zajmujących się sztuczną inteligencją (AI), blockchainem i technologiami zdecentralizowanymi, w tym FuixLabs i MiraScribe. Jego doświadczenie zawodowe obejmuje stanowiska starszego inżyniera w Google, gdzie zajmował się rozproszonymi systemami sieciowymi i transpilacją języków, a także stanowiska w firmach Hewlett-Packard i Cerner Corporation. O'Brien posiada bogate doświadczenie w budowaniu systemów o znaczeniu krytycznym, zarządzaniu łańcuchem dostaw oraz aplikacjach fintech. Jest również doświadczonym konsultantem, który pracował nad systemami inteligentnych kontraktów i zautomatyzowanym tworzeniem rynku. Obecnie koncentruje się na stosowaniu najlepszych praktyk inżynierii oprogramowania w rozwoju solidnych, skalowalnych i bezpiecznych generatywnych systemów AI i wieloagentowych.

Tim O'Brien is a technology leader, architect, and developer with over 20 years of experience in the software industry. He has served as a CTO and founder of multiple companies focused on AI, blockchain, and decentralized technologies, including FuixLabs and MiraScribe. His professional background includes roles as a Senior Engineer at Google, where he worked on distributed network systems and language transpilation, as well as positions at Hewlett-Packard and Cerner Corporation. O'Brien has extensive expertise in building mission-critical systems, supply chain management, and fintech applications. He is also an experienced consultant, having worked on smart contract systems and automated market-making. He currently focuses on applying software engineering best practices to the development of robust, scalable, and secure generative AI and multi-agent systems.

FuixLabs

Literatura uzupełniająca

Książki

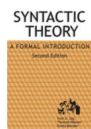
Literatura pokrewna (Google Scholar):

[1] "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? 🦜"

Emily M. Bender

[2] "Data and its (dis)contents: A survey of dataset development and use in machine learning research"

Emily M. Bender



[3] "Syntactic Theory: A Formal Introduction"

Emily M. Bender

Center for the Study of Language and Information Publication

[4] "Distributed Artificial Intelligence Research Institute"

Timnit Gebru

[5] "Black in AI"

Timnit Gebru

[6] "Návrh programů pomocí vzorů: stavební kameny objektově orientovaných programů"

Erich Gamma

ISBN: 9788024703022



[7] "Design Patterns"

Erich Gamma

Pearson Education | ISBN: 9780321700698



[8] "JUnit"

Erich Gamma

Simon and Schuster | ISBN: 9781638356554



[9] "Enterprise Integration Patterns"

Gregor Hohpe

Addison-Wesley | ISBN: 9780133065107

Artykuły naukowe

Autorzy przywoływani:

[1] "Message from Noam Chomsky"

Noam Chomsky

2006 | Lingua | DOI: 10.1016/j.lingua.2006.06.001

[Google Scholar](#)

[2] "Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data"

Emily M. Bender, Alexander Koller

2020 | DOI: 10.18653/v1/2020.acl-main.463

[Google Scholar](#)

[3] "Data Statements for Natural Language Processing: Toward Mitigating System Bias and Enabling Better Science"

Emily M. Bender, Batya Friedman

2018 | Transactions of the Association for Computational Linguistics | DOI: 10.1162/tacl_a_00041

[Google Scholar](#)

[4] "Data and its (dis)contents: A survey of dataset development and use in machine learning research"

Amandalynne Paullada, Inioluwa Deborah Raji, Emily M. Bender, Emily Denton, Alex Hanna

2020 | arXiv (Cornell University) | DOI: 10.1016/j.patter.2021.100336

[Google Scholar](#)

[5] "Situating Search"

Chirag Shah, Emily M. Bender

2022 | DOI: 10.1145/3498366.3505816

[Google Scholar](#)

[6] "Using deep learning and Google Street View to estimate the demographic makeup of neighborhoods across the United States"

Timnit Gebru, Jonathan Krause, Yilun Wang, Duyun Chen, Jia Deng

2017 | Proceedings of the National Academy of Sciences | DOI: 10.1073/pnas.1700035114

[Google Scholar](#)

[7] "Lessons from archives"

Eun Seo Jo, Timnit Gebru

2020 | DOI: 10.1145/3351095.3372829

[Google Scholar](#)

[8] "The TESCREAL bundle: Eugenics and the promise of utopia through artificial general intelligence"

Timnit Gebru, Émile P. Torres

2024 | First Monday | DOI: 10.5210/fm.v29i4.13636

[Google Scholar](#)

[9] "Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing"

Inioluwa Deborah Raji, Andrew Smart, Rebecca N. White, Margaret Mitchell, Timnit Gebru

2020 | arXiv (Cornell University) | DOI: 10.48550/arxiv.2001.00973

[Google Scholar](#)

[10] "Knowledge of language: its elements and origins"

Noam Chomsky

1981 | Philosophical transactions of the Royal Society of London. Series B, Biological sciences | DOI: 10.1098/rstb.1981.0135

[Google Scholar](#)

[11] "Editors' introduction: some concepts and issues in linguistic theory"

Noam Chomsky

2002 | Cambridge University Press eBooks | DOI: 10.1017/cbo9780511613876.002

[Google Scholar](#)

[12] "Design Patterns: Elements of Reusable Object-Oriented Software"

Erich Gamma, Richard F. Helm, Ralph E. Johnson, John Vlissides

1994

[Google Scholar](#)

[13] "Design Patterns: Abstraction and Reuse of Object-Oriented Design"

Erich Gamma, Richard F. Helm, Ralph E. Johnson, John Vlissides

1993 | Lecture notes in computer science | DOI: 10.1007/3-540-47910-4_21

[Google Scholar](#)

[14] "Contributing To Eclipse: Principles, Patterns, And Plug-Ins"

Erich Gamma, Kent Beck

2003

[Google Scholar](#)

[15] "ET++—an object oriented application framework in C++"

André Weinand, Erich Gamma, Rudolf Marty

1988 | DOI: 10.1145/62083.62089

[Google Scholar](#)

[16] "Patterns of optimized loops"

Samira Tasharofi, Ralph Johnson

2010 | Proceedings of the 2010 Workshop on Parallel Programming Patterns | DOI: 10.1145/1953611.1953617

[Google Scholar](#)

[17] "Patterns for cache optimizations on multi-processor machines"

Nicholas Chen, Ralph Johnson

2010 | Proceedings of the 2010 Workshop on Parallel Programming Patterns | DOI: 10.1145/1953611.1953613

[Google Scholar](#)

[18] "Components, frameworks, patterns"

Ralph E. Johnson

1997 | Proceedings of the 1997 symposium on Software reusability - SSR '97 | DOI: 10.1145/258366.258378

[Google Scholar](#)

[19] "Patterns for evolving frameworks"

Don Roberts, Ralph E. Johnson

1997

[Google Scholar](#)

[20] "Pattern languages of program design"

John Vlissides, James O. Coplien, Norman L. Kerth

1995 | Medical Entomology and Zoology

[Google Scholar](#)

[21] "Automatic code generation from design patterns"

Frank Budinsky, M. A. Finnie, John Vlissides, P.S. Yu

1996 | IBM Systems Journal | DOI: 10.1147/sj.352.0151

[Google Scholar](#)

[22] "Enterprise Integration Patterns: Designing, Building, and Deploying Messaging Solutions"

Gregor Hohpe, Bobby Woolf

2003

[Google Scholar](#)

[23] "Enterprise Integration Patterns"

Gregor Hohpe, Bobby Woolf

2003

[Google Scholar](#)

[24] "The Software Architect's Role in the Digital Age"

Gregor Hohpe, İpek Özkaya, Uwe Zdun, Olaf Zimmermann

2016 | IEEE Software | DOI: 10.1109/ms.2016.137

[Google Scholar](#)

[25] "A Decade of Enterprise Integration Patterns: A Conversation with the Authors"

Olaf Zimmermann, Cesare Pautasso, Gregor Hohpe, Bobby Woolf

2015 | IEEE Software | DOI: 10.1109/ms.2016.11

[Google Scholar](#)

[26] "Mobile IP : design principles and practices"

Charles E. Perkins, Sherman R. Alpert, Bobby Woolf

1997 | Medical Entomology and Zoology

[Google Scholar](#)

[27] "The design patterns Smalltalk companion"

Sherman R. Alpert, Kyle Brown, Bobby Woolf

1998 | Medical Entomology and Zoology

[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: c150ce84-43f7-4a62-8192-2b3d1c766f6d