



Fundusze Europejskie
dla Rozwoju Społecznego



Rzeczpospolita
Polska

Dofinansowane przez
Unię Europejską



Politechnika Świętokrzyska
Wydział Zarządzania i Modelowania Komputerowego

Kierunek studiów:
Inżynieria danych

Dariusz Dobrowolski

Materiały dydaktyczne do przedmiotu

GRAFOWE BAZY DANYCH

opracowane w ramach realizacji Projektu
**„Dostosowanie kształcenia
w Politechnice Świętokrzyskiej do potrzeb
współczesnej gospodarki”**
FERS.01.05-IP.08-0234/23

Kielce, 2025



Politechnika Świętokrzyska
Kielce University of Technology

Projekt „Dostosowanie kształcenia w Politechnice Świętokrzyskiej do potrzeb współczesnej gospodarki”
nr FERS.01.05-IP.08-0234/23



Spis treści

1. Wprowadzenie	4
2 Podstawy grafowych baz danych.....	6
2.1 Definicja grafowych baz danych.....	6
2.2 Podstawowe pojęcia teorii grafów.....	8
2.3 Architektura systemów grafowych.....	14
3 Historia rozwoju baz danych	16
3.1 Ewolucja technologii bazodanowych.....	16
3.2 NoSQL i przejście do grafów	20
3.3 Wpływ rozwoju technologii na zastosowania grafów.....	22
4 Modelowanie danych w bazach grafowych	25
4.1 Reprezentacja danych w postaci grafu	25
4.2 Projektowanie schematu grafu.....	27
5 Języki zapytań w bazach grafowych	31
5.1 Przegląd języków zapytań	31
5.2 Cypher	33
5.3 Gremlin	37
6 Wydajność i skalowalność baz grafowych	42
6.1 Optymalizacja zapytań.....	42
6.2 Przechowywanie i indeksowanie danych	44
7 Zastosowania baz grafowych.....	46
7.1 Sieci społecznościowe	46
7.2 Systemy rekomendacyjne	50
7.3 Analiza sieci powiązań biznesowych	53
7.4 Bezpieczeństwo i wykrywanie oszustw.....	55
8 Integracja baz grafowych z innymi technologiami.....	59
8.1 Integracja z relacyjnymi bazami danych	59





8.2 Integracja z systemami NoSQL	61
9 Przyszłość baz grafowych.....	63
9.1 Nowe kierunki badań	63
9.2 Wpływ sztucznej inteligencji na rozwój baz grafowych	65
10 Zakończenie	68
11. Literatura.....	70



Materiały dydaktyczne objęte licencją Creative Commons BY 4.0.

Licencja dostępna pod adresem: <https://creativecommons.org/licenses/by/4.0/>



1. Wprowadzenie

Grafowe bazy danych bazują na idei przechowywania informacji jako zbioru wierzchołków oraz krawędzi reprezentujących relacje między nimi. Taki sposób organizacji pozwala odzwierciedlić złożone, często gęsto powiązane struktury, jakie spotyka się w wielu aspektach analizy danych. W odróżnieniu od tradycyjnych systemów opartych na modelu relacyjnym, gdzie informacje są dzielone na tabele powiązane kluczami obcymi, w grafowym podejściu relacje są elementami pierwszoplanowymi i dostęp do nich odbywa się bez konieczności wykonywania kosztownych operacji złączeń. Pozwala to skrócić czas przetwarzania zapytań, zwłaszcza przy danych o dużej liczbie połączeń.

Architektura baz grafowych, takich jak Neo4j, uwzględnia natywne mechanizmy przechowywania grafów, co oznacza, że ich struktura jest projektowana od podstaw w taki sposób, by optymalizować operacje przechodzenia przez węzły i krawędzie. Przechowywanie informacji w formacie grafowym umożliwia realizację zapytań obejmujących na przykład wyszukiwanie ścieżek czy analizę powiązań między elementami bez utraty wydajności wynikającej z przekształceń danych do innych modeli. Tego typu rozwiązanie nie jest dodatkiem do istniejącej technologii, ale integralnym fundamentem systemu. Istnieje różnorodność zastosowań baz grafowych. W obszarach takich jak sieci społecznościowe krawędzie mogą reprezentować znajomości lub interakcje pomiędzy użytkownikami. Dzięki temu możliwe jest szybkie określenie stopnia powiązań czy rekomendowanie nowych kontaktów na podstawie struktury całego grafu. Systemy rekomendacyjne wykorzystują podobną logikę, gdzie relacje łączą użytkowników z produktami i ich ocenami. Analiza takiego grafu pozwala efektywnie przewidywać preferencje. W kontekście biznesowym bazy grafowe znajdują zastosowanie także w zarządzaniu danymi hierarchicznymi i infrastrukturalnymi. Przykładem może być modelowanie struktur organizacyjnych czy sieci telekomunikacyjnych zawierających zależności pomiędzy planami taryfowymi, klientami i urządzeniami.

W przypadkach, kiedy dane są silnie powiązane i podlegają częstym aktualizacjom oraz zapytaniom ad hoc, tradycyjne rozwiązania relacyjne często okazują się niewydajne ze względu na konieczność stosowania licznych złączeń. Neo4j, będący jednym z najpopularniejszych systemów tego typu, wyróżnia się dzięki językowi zapytań Cypher. Jest on deklaratywny i inspirowany składnią SQL, co ułatwia jego stosowanie dla osób znających technologie bazodanowe oparte na językach zapytań. Cypher umożliwia intuicyjne definiowanie wzorców wyszukiwanych w grafie oraz filtrowanie wyników zgodnie z zadanymi ograniczeniami. To sprawia, że interakcja z danymi staje się naturalna zarówno dla programistów, jak i analityków.



Warto zwrócić uwagę na istotne różnice pomiędzy architekturą natywną a nienatywną w odniesieniu do przetwarzania grafowego. Technologie określane jako "graph first" projektowane są do obsługi danych grafowych od najniższego poziomu i implementują mechanizmy ACID dla bezpieczeństwa transakcji. Z kolei systemy nienatywne dodają obsługę grafu jedynie jako rozszerzenie istniejących struktur magazynowania innego rodzaju danych. Skutkuje to często mniejszą wydajnością podczas intensywnego korzystania z relacji między obiektami. Tradycyjne modele agregatowe wykorzystywane w wielu typach baz NoSQL – takie jak magazyny klucz-wartość czy bazy dokumentów – unikają wyraźnego odwzorowywania połączeń między obiektami. To sprawia, że analiza zależności wymaga dodatkowych etapów obliczeniowych lub wczytywania wielu fragmentów danych naraz.

Bazy grafowe eliminują ten problem dzięki możliwości bezpośredniego przemieszczania się po relacjach zapisanych jako krawędzie. Rozważając aspekt praktyczny wdrożeń warto przytoczyć przypadki firm przerzucających swoje procesy z baz relacyjnych właśnie na systemy grafowe. Na przykład Cisco zastosowało Neo4j do zarządzania hierarchicznymi strukturami produktów oraz danymi użytkowników, uzyskując dzięki temu natychmiastowy dostęp do informacji zamiast czasochłonnych zapytań wykonywanych dotąd w Oracle RAC. Portale społecznościowe takie jak Facebook czy LinkedIn przechowują swoje ogromne sieci kontaktów właśnie jako grafy. Kwestia wydajności ma tu zasadnicze znaczenie. Analiza ścieżek najkrótszych czy wyszukiwanie klastrów wymaga algorytmów działających iteracyjnie i często ingerujących bezpośrednio w strukturę danych. Baza grafowa umożliwia realizację tych scenariuszy bez nadmiernej komplikacji technologicznej – przykładowo wyszukanie spójnych komponentów wewnątrz dużej sieci opisanej jako zbiór węzłów następuje poprzez sekwencyjne przechodzenie krawędzi i zaznaczanie odwiedzonych elementów.

Nie można jednak pominąć kwestii dopasowania konkretnego rozwiązania do określonych wymagań projektu. Analizy porównawcze wskazują różnice pomiędzy popularnymi implementacjami jak Neo4j, GraphDB czy OrientDB zarówno pod kątem typowych operacji CRUD jak i wydajności przy różnych rodzajach zapytań. Każdy system posiada zestaw cech determinujących jego przydatność w określonym kontekście – od kompatybilności językowej po możliwości optymalizacji pamięci masowej, co wiąże się bezpośrednio z architekturą jego natywnego magazynu. Tak skonstruowane środowisko pracy z danymi daje możliwość nie tylko szybkiej odpowiedzi na pytania o strukturę połączeń, lecz również łatwego rozbudowywania modelu o nowe typy relacji czy atrybuty bez przebudowy całej bazy. Możliwość adaptacji schematu do zmieniających się warunków biznesowych stanowi jeden z filarów rosnącej popularności technologii baz grafowych.



2 Podstawy grafowych baz danych

2.1 Definicja grafowych baz danych

Grafowa baza danych jest systemem, w którym fundamentalnym sposobem organizacji informacji jest model grafu złożonego z wierzchołków i krawędzi. Wierzchołki reprezentują jednostki lub obiekty, takie jak użytkownicy, produkty czy dokumenty, podczas gdy krawędzie opisują powiązania pomiędzy tymi jednostkami. Relacje te są elementem pierwszoplanowym i istnieją w systemie jako odrębne struktury o własnych właściwościach. Model ten różni się zasadniczo od baz relacyjnych, gdzie więzi są realizowane poprzez klucze obce i wymagają wykonywania operacji złączeń, co często powoduje wzrost kosztów obliczeniowych przy danych o wysokim stopniu powiązań.

Istotną cechą grafowych baz danych jest ich zdolność do bezpośredniego przechodzenia pomiędzy powiązanymi obiektami dzięki mechanizmowi tzw. przetwarzania natywnego. W tym przypadku każdy wierzchołek posiada bezpośrednie odniesienia do swoich sąsiadów, co określa się mianem sąsiedztwo bez indeksu (bez indeksowe) (ang. index-free adjacency). Rozwiązanie takie eliminuje konieczność korzystania z indeksów globalnych przy nawigowaniu po strukturze grafu i przekłada się na wysoką wydajność operacji eksplorujących powiązania, zwłaszcza przy głębokich zapytaniach obejmujących wiele stopni relacji. Model grafowy jest elastyczny w zakresie struktury – umożliwia dodawanie nowych typów wierzchołków oraz krawędzi bez konieczności przebudowy całej bazy czy migracji istniejących danych. Jest to szczególnie widoczne w zastosowaniach, gdzie dane mają charakter dynamiczny i wymagają ciągłego rozszerzania modelu. Przykładowo w sieci społecznej możemy zacząć od prostego zbioru relacji typu „znajomy” pomiędzy użytkownikami, a następnie wzbogacić go o inne rodzaje kontaktów, jak „kolega z pracy” czy „obserwowany”, przy zachowaniu spójności całości.

W kontekście definicji należy podkreślić rozróżnienie między architekturą natywną a nienatywną. W pierwszym przypadku każdy element systemu – od języka zapytań po sposób zapisu danych na dysku – dostosowany jest specjalnie do obsługi grafów. W natywnym systemie, takim jak Neo4j, dane są przechowywane w dedykowanych plikach odpowiadających za konkretne komponenty: osobno dla wierzchołków, relacji oraz etykiet czy właściwości. Krawędzie są tu pełnoprawnymi rekordami zawierającymi wskaźniki do węzłów źródłowego i docelowego oraz atrybuty opisowe. Wersje nienatywne opierają się na ogólnych magazynach danych (relacyjnych lub kolumnowych), a warstwa obsługi grafu jest dodana jako nadbudowa – często ze stratą wydajności przy intensywnej analizie relacji.



Grafowa baza danych nie ogranicza się jednak do samego modelowania. Istotny jest także sposób interakcji z danymi za pomocą dedykowanego języka zapytań. Cypher używany w Neo4j został opracowany tak, aby intuicyjnie odwzorować model grafowy – jego składnia opiera się na symbolice kół jako węzłów oraz strzałek jako krawędzi łączących je. Pozwala to wyrażać zapytania jako wzorce struktur, które system ma odnaleźć w bazie. Zaletą jest spójność pomiędzy tym, jak projektujemy model, a jak później formułujemy pytania badawcze wobec danych. Problematyka kosztów operacji stanowi część definicji praktycznej. Ponieważ wiele zastosowań dotyczy grafów o dużych rozmiarach i wysokiej dynamice zmian – przykładem mogą być logi zdarzeń czy interakcje użytkowników generujące tysiące połączeń na sekundę – ważne jest wsparcie zarówno dla zapytań punktowych, jak i analiz globalnych obejmujących cały graf lub jego fragmenty. System musi umożliwiać efektywne wydobywanie subgrafów z większej struktury oraz przechowywanie wersji historycznych pozwalających śledzić ewolucję relacji.

Na tle innych magazynów danych NoSQL warto zaznaczyć różnicę podejścia do relacji. Magazyny klucz-wartość czy dokumentowe mogą przechowywać identyfikatory innych obiektów jako pola wewnętrzne danego rekordu, co przypomina użycie kluczy obcych. Jednak takie podejście wymaga wykonywania kosztownych operacji łączenia na poziomie aplikacyjnym oraz utrudnia realizację zapytań dwukierunkowych bez dodatkowych zabiegów (jak dodawanie relacji odwrotnych). Bazy grafowe unikają tego problemu przez umieszczenie połączeń w samym centrum modelu – relacje są zasobami pierwszego rzędu i posiadają takie samo znaczenie jak same dane.

Definicyjny obraz bazy grafowej obejmuje także możliwość stosowania silników obliczeniowych dedykowanych analizom grafowym. Tak zwane grafowe silniki obliczeniowe pozwalają uruchamiać algorytmy analizujące cały zbiór danych pod kątem właściwości takich jak klastrowanie czy średnia liczba połączeń na wierzchołek. Choć część z nich zawiera własną warstwę magazynowania, wiele działa na bazach statycznych lub strumieniach zasilanych spoza systemu bazodanowego. Ostatecznie grafowa baza danych to nie tylko schemat czy technologia przechowywania – to również filozofia pracy z informacją skoncentrowaną na zależnościach między bytami. Tak rozumiana definicja obejmuje zarówno formalny opis struktury (węzły plus krawędzie), sposób zapisu i przetwarzania (natywnie versus zaadaptowane technologie), mechanizmy wyszukiwania (języki wzorcowe typu Cypher) oraz kontekst zastosowań wymagających dynamicznego modelowania powiązań bez naruszenia istniejącej integralności danych.





2.2 Podstawowe pojęcia teorii grafów

Wierzchołki i krawędzie

W modelu grafowym wierzchołki oraz krawędzie stanowią fundamentalne elementy struktury danych, definiujące sposób przechowywania i organizacji informacji. Wierzchołek (nazywany także węzłem) jest jednostką reprezentującą obiekt lub byt w analizowanym systemie. Może to być osoba, produkt, dokument, serwer czy dowolny inny element opisywany przez zbiór właściwości. Właściwości odpowiadają atrybutom o charakterze klucz-wartość, gdzie klucz jest ciągiem znaków opisującym daną cechę, a wartość może przyjmować praktycznie dowolny typ obsługiwany przez bazę. Krawędzie natomiast służą odzwierciedlaniu relacji pomiędzy parami wierzchołków – mogą mieć kierunek lub pozostawać nieukierunkowane, zależnie od potrzeb modelu i semantyki danej relacji. Każda krawędź posiada etykietę będącą nazwą typu relacji. Etykieta pełni funkcję semantyczną, informując o tym, w jaki sposób powiązane są ze sobą dwa konkretne wierzchołki. Oprócz samej nazwy krawędzi możliwe jest przypisanie jej zestawu właściwości analogicznych do tych występujących w wierzchołkach. Umożliwia to zapis dodatkowych informacji kontekstowych – na przykład daty utworzenia połączenia pomiędzy użytkownikami lub wartości liczbowej określającej intensywność danej interakcji.

Struktura grafu oparta na takich komponentach pozwala zachować spójność danych nawet wtedy, gdy modelowanie obejmuje wielowymiarowe i złożone powiązania. W przypadku tzw. binarnych relacji (gdzie połączenie dotyczy dokładnie dwóch bytów) naturalnym sposobem odwzorowania ich jest użycie krawędzi łączącej odpowiednie wierzchołki. Dla relacji n-arnych (więcej niż dwa uczestniczące obiekty) praktycznym rozwiązaniem bywa traktowanie samej relacji jako odrębnego wierzchołka połączonego krawędziami z każdym z elementów biorących udział w powiązaniu. Pozwala to uniknąć komplikacji związanych z nadmiarem metadanych w jednej krawędzi oraz ułatwia rozbudowę modelu o kolejne elementy relacji.

W wielu implementacjach baz grafowych istotne jest jednoznaczne etykietowanie zarówno wierzchołków jak i krawędzi, aby aplikacja mogła szybciej identyfikować ich typ oraz przeznaczenie. W nowoczesnych wersjach silników takich jak Tinkerpop czy Neo4j etykiety stały się obowiązkowym elementem definicji struktury – dzięki temu możliwe jest stosowanie indeksów czy ograniczeń dla określonych klas obiektów. Label (etykieta) może pełnić funkcję typu z perspektywy aplikacyjnej: przykładowo etykieta „User” przypisana do wierzchołka oznacza typ użytkownika systemu. Konstrukcja grafu zgodna z modelem właściwości (property graph) zakłada istnienie zbioru V reprezentującego wszystkie wierzchołki oraz zbioru E zawierającego wszystkie. Formalnie graf można zapisać jako



$G=(V,E)$

gdzie każdy element $e \in E$ wiąże uporządkowaną parę (v_out , v_in) będącą wskaźnikami do konkretnego źródła i celu relacji. Index-free adjacency oznacza, że v_out i v_in bezpośrednio wskazują na konkretne rekordy w pamięci lub na dysku, co przyspiesza eksplorację grafu podczas zapytań.

Kwestia kierunkowości relacji wymaga szczególnej uwagi przy projektowaniu schematu danych. Przyjęcie kierunku umożliwia odróżnienie roli poszczególnych uczestników połączenia – na przykład „User” → „LIKES” → „Product” sugeruje preferencje użytkownika względem produktu, podczas gdy odwrotna relacja mogłaby oznaczać rekomendację ze strony produktu dla użytkownika. Relacje bezkierunkowe sprawdzają się tam, gdzie semantyka łączy oba obiekty równoważnie – np. „isSiblingOf” pomiędzy dwoma osobami.

W praktyce operacje na grafach wymagają jasnego wzorcowego opisu struktury zapytania. Cypher stosowany m.in. przez Neo4j wykorzystuje czytelną notację graficzną przypominającą diagramy – nawiasy okrągłe ‘()’ symbolizują wierzchołki wraz z opcjonalnymi właściwościami zapisanymi wewnątrz nawiasów klamrowych “, zaś strzałki ‘->’ wskazują kierunek powiązań. Przykład ‘(a:Person name:’Jim’)-[:KNOWS]->(b)’ pokazuje szukanie osoby o imieniu Jim posiadającej znajomego b.

Projektując system należy też rozważyć optymalizację poprzez kontrolę liczby właściwości przypisanych bezpośrednio do krawędzi i wierzchołków – nadmiar szczegółów może zwiększyć koszty pamięciowe i utrudnić indeksowanie. Z drugiej strony brak kluczowych metadanych ogranicza możliwości analizy kontekstowej, zwłaszcza we wdrożeniach obejmujących dane dynamiczne jak sieci społecznościowe czy analizy logów zdarzeń. Ważne jest także dostosowanie schematu etykiet do logiki aplikacyjnej: niektóre systemy wykorzystują etykiety wyłącznie do celów technicznych (np. segmentacja danych na potrzeby różnych serwisów), inne traktują je jako część modelu domenowego opisując rzeczywiste typy obiektów i powiązań. Oba podejścia mają swoje zalety – pierwsze upraszcza zarządzanie strukturą podczas migracji lub integracji danych między różnymi źródłami, drugie wspiera komunikację między zespołem programistycznym a analitykami domenowymi. Ostatecznie rola wierzchołków i krawędzi sprowadza się do zapewnienia możliwie naturalnego odwzorowania struktury logicznej problemu badawczego. Każdy scenariusz analizy determinuje jakie typy połączeń powinny znaleźć odzwierciedlenie w modelu oraz jakie informacje muszą być dostępne bezpośrednio przy przeglądaniu powiązań. Dobrze zaprojektowana para V–E pozwala tworzyć zarówno płytkie zapytania punktowe jak i głębokie operacje obejmujące całe podgrafy bez znaczącego spadku wydajności.



Atrybuty i właściwości

W modelu grafowym atrybuty i właściwości stanowią podstawowy mechanizm wzbogacania semantyki danych o kontekst opisowy. Każdy wierzchołek może posiadać zestaw par klucz–wartość, gdzie klucz pełni rolę identyfikatora cechy, a wartość jej konkretnego opisu liczbowego, tekstowego lub dowolnego innego obsługiwanego typu. Takie właściwości pozwalają na przechowywanie metadanych związanych bezpośrednio z obiektem – mogą to być imiona i nazwiska użytkowników, daty utworzenia wpisów czy adresy URL zasobów. Analogiczny mechanizm funkcjonuje również dla krawędzi, co umożliwia zapis dodatkowych informacji dotyczących relacji, np. wagi połączenia, kosztu transakcji czy odległości geograficznej między lokalizacjami. Właściwości przypisane do krawędzi mają szczególne znaczenie w analizach wymagających uwzględnienia wartości liczbowych określających charakter relacji. Jeżeli relacje opisują czasy odpowiedzi serwera dla danego użytkownika, ich właściwości mogą służyć do filtrowania bądź sortowania wyników, bez konieczności przechodzenia przez cały zbiór danych.

Model grafowy zapewnia traktowanie takich detali jako integralnej części struktury połączeń. Ponieważ zarówno wierzchołki, jak i krawędzie posiadają unikalne identyfikatory oraz mogą występować w dowolnej liczbie i typie relacji, system jest w stanie przechowywać różnorodne zestawy atrybutów odpowiadające wymaganiom domeny aplikacyjnej. Praktyka projektowa pokazuje, że dobór właściwości wymaga równowagi między zakresem informacji a kosztami ich utrzymania. Nadmiar danych szczegółowych może komplikować indeksowanie oraz zwiększać rozmiary plików przechowujących właściwości w magazynie systemowym – Neo4j wykorzystuje tu odrębny plik *neostore.propertystore.db*, gdzie właściwości zapisane są jako lista pojedynczo połączonych rekordów. Każdy rekord zawiera ustaloną liczbę bloków, co pozwala przewidzieć fizyczną organizację danych na dysku i zoptymalizować dostęp, ale zarazem ogranicza elastyczność zapisu bardzo dużych zestawów wartości przypisanych do jednego elementu. Istotnym aspektem jest także rola etykiet i typów jako mechanizmu grupowania zbiorów właściwości. Etykieta wskazuje kategorię obiektu lub relacji – przykładowo etykieta "Person" może odpowiadać zestawowi właściwości takich jak name, birthDate, email. Kiedy ten sam typ obiektu pojawia się wielokrotnie z różnymi wartościami atrybutów, etykiety pomagają zachować spójność przetwarzania oraz umożliwiają stosowanie indeksów lokalnych dla właściwości wybranego typu.

W kontekście zapytań istotnym mechanizmem jest punktowe wyszukiwanie według wartości właściwości, które polega na wskazaniu konkretnej wartości dla danej cechy i odnalezieniu obiektów ją posiadających. Ze względu na strukturę sąsiedztwa bez indeksu można takie zapytanie realizować szybko przy niewielkich głębokościach



przechodzenia po grafie; natomiast przy bardziej złożonych filtrach obejmujących jednocześnie warunki dla wielu właściwości konieczne staje się wsparcie indeksowania dedykowanego. Relacje mogą być wzbogacane o parametry jakościowe czy ilościowe nie tylko z myślą o filtrowaniu wyników zapytań, ale również do celów analitycznych korzystających z modeli matematycznych.

W analizach przepływu sieciowego krawędzie mogą przechowywać atrybut `capacity` używany w algorytmach maksymalnego przepływu. W sieciach społecznościowych wartość `interactionFrequency` może posłużyć do wyznaczania siły więzi. Rozumienie roli takich właściwości pozwala lepiej projektować logikę zapytań oraz algorytmy działające w obrębie bazy grafowej. Warto zwrócić uwagę na kompatybilność różnych modeli grafowych – grafy właściwości domeny rozszerzają możliwości klasycznych grafów właściwości poprzez umieszczanie adnotacji poza samą strukturą grafu. Takie podejście umożliwia modyfikację interpretacji wartości właściwości bez ingerencji w bazową strukturę wierzchołków i krawędzi oraz ułatwia współpracę z systemami dziedziczenia przechowującymi starsze formaty własności (np. napisy określające płeć czy etykiety kategorii jako oddzielne węzły). Z perspektywy implementacyjnej przypisanie właściwości powinno uwzględniać możliwe ograniczenia sprzętowe i architektoniczne systemu bazodanowego.

Niektóre silniki optymalizują dostęp do często używanych atrybutów poprzez przechowywanie ich w pamięci podręcznej lub zastosowanie kodowania binarnego skracającego czas deserializacji. Inne rezygnują z takiej optymalizacji na rzecz prostoty modelu kosztem wydajności przy intensywnym wykorzystaniu wielu typów właściwości. Ostatecznie rola atrybutów i właściwości sprowadza się do wzbogacenia czystej topologii grafu o dane definiujące jego kontekst funkcjonalny i semantyczny. Dzięki nim analiza przestaje być abstrakcyjnym badaniem układu połączeń i staje się narzędziem pozwalającym na powiązanie struktury z realnymi cechami obiektów i relacji. Umiejętne zaprojektowanie ich zakresu oraz sposobu przechowywania jest krytyczne dla osiągnięcia dobrej równowagi pomiędzy elastycznością modelowania a wydajnością operacyjną całego systemu grafowego.

Rodzaje grafów

W kontekście grafowych baz danych wyróżnia się kilka typów struktur grafowych, z których każda odpowiada na inne potrzeby modelowania i analizy danych. Najbardziej znanym jest graf właściwości, w którym zarówno wierzchołki, jak i krawędzie mogą posiadać zestawy par klucz–wartość opisujących ich cechy. Ten model daje dużą swobodę semantycznego opisywania obiektów oraz ich relacji, a także łatwość rozszerzania zakresu informacji bez modyfikacji samej topologii grafu. Etykiety przypisywane do wierzchołków oraz typy krawędzi wspierają logiczne





grupowanie elementów, co przekłada się na efektywne filtrowanie i indeksowanie. W ramach tego modelu można wyróżnić odmianę nazwaną wykresem domeny właściwości, gdzie wartości właściwości są reprezentowane jako dodatkowe połączenia, na przykład (n1,"gender","female"), pozwalając odwzorować atrybuty jako oddzielne elementy struktury. To podejście wspiera kompatybilność ze źródłami danych o odmiennym schemacie przechowywania metadanych.

Innym typem jest graf RDF (Resource Description Framework), który bazuje na trójkach postaci (x, y, z), znanych jako triple store. W tym modelu każdy fakt zapisany jest niezależnie w formacie podmiot–predykat–obiekt. Pozwala to na łatwą integrację z danymi opisowymi w standardzie semantycznego Internetu, choć jednocześnie utrudnia szybkie przechodzenie relacji ze względu na konieczność odtwarzania połączeń z wielu osobnych rekordów. Triple stores sprawdzają się raczej w analizach offline, gdzie większe znaczenie ma skalowalność horyzontalna niż minimalizacja czasu odpowiedzi systemu.

Istnieją także implementacje bazujące na hiper grafach (hypergraphs). W odróżnieniu od klasycznych grafów binarnych umożliwiają one definiowanie relacji wieloelementowych bez redukcji ich do zbioru par połączonych krawędzi. Hiper graf może być więc używany do uchwycenia złożonych zależności kontekstowych lub meta-intencji – na przykład reprezentowania transakcji obejmującej jednocześnie kilka obiektów powiązanych wspólnym wydarzeniem. Z punktu widzenia wydajności baza hipergrafowa wymaga innych metod przeszukiwania niż graf właściwości, co może komplikować standardowe algorytmy eksploracyjne. Model wykresu domeny stosowany m.in. w MillenniumDB formalizuje zapis krawędzi jako relacji w postaci "DomainGraph(source,type,target,eid)", gdzie eid jest identyfikatorem krawędzi. Tak skonstruowany graf podporządkowuje model RDF, ale oferuje większą elastyczność w rozszerzaniu typów połączeń czy dodawaniu metadanych poprzez mechanizmy takie jak edge-as-node czy nested edge nodes. Edge-as-node umożliwia traktowanie krawędzi jako pełnoprawnego wierzchołka i łączenie jej z innymi elementami, co poszerza możliwości analizy zależności o charakterze meta poziomym – np. łączenie relacji „zakup” z jednostką opisującą konkretną kampanię marketingową. W praktyce spotyka się także konstrukcje specjalistyczne takie jak tzw. wirtualne krawędzie dostępne m.in. w Sparksee.

Tego typu krawędzie nie są materializowane w magazynie danych – istnieją logicznie poprzez regułę łączenia wierzchołków o tej samej wartości danego atrybutu. Krawędź wirtualna pozwala uprościć odwzorowanie pewnych wzorców bez fizycznego powielania relacji, jednak wiąże się to z koniecznością ich generowania podczas zapytań, co może wpływać na czas odpowiedzi.



Z perspektywy zastosowań rozróżnienie pomiędzy grafami ukierunkowanymi oraz nieukierunkowanymi ma istotne znaczenie dla interpretacji danych. Graf ukierunkowany definiuje kierunek przepływu informacji lub zależności pomiędzy elementami – przykładowo strzałka od klienta do zamówienia wskazuje inicjatora zdarzenia. Graf nieukierunkowany oznacza symetrię relacji, często spotykaną np. przy modelowaniu przyjaźni czy podobieństwa produktów. Odrębne miejsce zajmują też multigrafy, dopuszczające istnienie wielu krawędzi łączących tę samą parę wierzchołków – przydatne np. do rejestrowania różnych kanałów komunikacji pomiędzy dwoma osobami.

W konstruowaniu systemu należy uwzględniać również możliwość obsługi grafów z etykietami hierarchicznymi lub kompozycyjnymi. Niektóre implementacje pozwalają przypisać więcej niż jedną etykietę do elementu oraz traktować je jako klasę zbiorową, taki mechanizm umożliwia przesuwanie obiektu pomiędzy kontekstami analitycznymi bez utraty jego dotychczasowych właściwości. Z kolei modele mieszane starają się łączyć zalety kilku typów grafów jednocześnie. Można spotkać rozwiązania integrujące grafy właściwości z mechaniką RDF tak, by część danych zachowywała wysoką szybkość przetwarzania dzięki strukturze właściwościowej, a inna była zgodna ze standardami wymiany informacji sieci semantycznej. Analogicznie niektóre platformy oferują tryb przetwarzania natywnego dla części grafu oraz warstwę translacyjną dla integracji z backendem relacyjnym czy dokumentowym.

Pod kątem analitycznym wybór rodzaju grafu wpływa także na dostępność określonych algorytmów optymalizacji lub eksploracji danych. System obsługujący sąsiedztwo bez indeksowe dla ukierunkowanego grafu właściwości będzie wyraźnie szybszy przy wyszukiwaniu ścieżek niż triple store wymagający rekonstrukcji powiązań ze zbioru niezależnych faktów. W przypadku hiper grafów konieczne bywa stosowanie bardziej zaawansowanych technik rozpisywania n-arnej relacji na zestaw operacji możliwych do wykonania przez silnik obliczeniowy. Szczególne warianty, takie jak zagnieżdżone węzły krawędziowe, rozszerzają możliwość budowania wielopoziomowych modeli relacyjnych poprzez iteracyjne traktowanie relacji jako bytów zdolnych do tworzenia kolejnych relacji. W środowiskach analitycznych pozwala to opisać proces wieloetapowej współpracy czy zależności cykliczne bez konieczności utrzymywania rozbudowanych struktur tabelarycznych.

Ostateczny dobór rodzaju grafu powinien uwzględniać zarówno aspekt wydajnościowy jak i zgodność ze źródłami danych oraz wymaganiami aplikacyjnymi, inne preferencje pojawią się przy budowie systemu rekomendacyjnego skupionego na analizie ścieżek użytkowników, a inne przy integracji rozproszonych repozytoriów wiedzy zgodnych z RDF i SPARQL.



2.3 Architektura systemów grafowych

Architektura systemów grafowych obejmuje zarówno warstwę przechowywania danych, jak i mechanizmy ich przetwarzania oraz interfejsy umożliwiające interakcję z zasobami. Rdzeń takiego systemu może być oparty na natywnym magazynie grafowym lub korzystać z magazynów nienatywnych, które wymagają dodatkowej warstwy translacyjnej do odwzorowania struktury grafu. Natywne przechowywanie cechuje się tym, że wierzchołki i krawędzie są zapisywane w dedykowanych strukturach pamięci masowej. Każdy element posiada unikalny identyfikator i jest powiązany bezpośrednimi wskaźnikami z innymi elementami, co umożliwia mechanizm sąsiedztwa bez indeksowego. Dzięki temu system zachowuje wysoką wydajność przy przechodzeniu po relacjach, nawet w przypadku głębokich zapytań.

Architektury nienatywne bazują najczęściej na relacyjnych lub obiektowych magazynach danych, w których graf jest serializowany do postaci tabel czy dokumentów. Ta metoda bywa wygodna przy integracji ze starszymi systemami lub narzędziami analitycznymi opartymi na SQL, jednak wprowadza dodatkowe opóźnienia podczas eksploracji powiązań, wynikające z konieczności odtwarzania struktury grafu z zapisów w formacie niegrafowym. Czasem stosuje się architekturę hybrydową – część danych przechowywana jest natywnie, a część w magazynie ogólnego przeznaczenia – co pozwala pogodzić wymagania dotyczące szybkości przetwarzania z potrzebami interoperacyjności. Warstwa przetwarzająca dane obejmuje silnik wykonujący operacje CRUD na modelu grafowym. W przypadku natywnego silnika kluczową rolę odgrywa metoda fizycznego „wskazywania” powiązanych elementów – relacje są reprezentowane przez parę wskaźników do wierzchołka źródłowego i docelowego wraz z atrybutami opisującymi połączenie. Taki zapis ułatwia stosowanie algorytmów eksploracyjnych, które iterują po krawędziach bez konieczności wyszukiwania adresów docelowych poprzez globalne indeksy. Silniki nienatywne obsługują grafy poprzez emulację tej logiki na bazie innych modeli danych; przykładowo przeszukanie sąsiadów danego wierzchołka wymaga wykonania zapytania filtrującego rekordy odpowiadające krawędziom powiązanym.

Jednym ze składników architektury jest także model danych, który determinuje sposób interpretacji przechowywanej struktury. Najpowszechniejszy graf właściwości wyróżnia zbiór V dla wierzchołków oraz E dla krawędzi, przy czym każdy element obu zbiorów może mieć własny zbiór właściwości.

Alternatywą jest model RDF operujący na trójkach podmiot–predykat–obiekt oraz bardziej zaawansowane grafy domenowe oferujące kompatybilność między różnymi formatami reprezentacji grafu poprzez wspólny rdzeń semantyczny. Dobór modelu



ma konsekwencje dla implementacji silnika – RDF wymaga obsługi wielu osobnych rekordów tworzących relację, podczas gdy graf właściwości łączy całość informacji o połączeniu w jednym obiekcie. Architekturę można rozszerzyć o komponenty odpowiedzialne za inferencję semantyczną, jak ma to miejsce w GraphDB.

Mechanizm ten generuje nowe fakty na podstawie istniejących poprzez zastosowanie reguł logicznych. Takie podejście integruje funkcje zarządzania wiedzą bez potrzeby ręcznej aktualizacji bazy przez użytkownika i może działać zarówno w trybie online (na żywo odpowiadając na zapytania), jak i jako proces cykliczny uruchamiany według harmonogramu.

W dużych instalacjach wydajność zależy także od architektury klastrowej i sposobu rozdzielania obciążenia pomiędzy serwerami. Systemy zgodne z ACID stosują często mechanizmy klastrowania master-slave zapewniające spójność transakcji kosztem skalowalności zapisu. Inne rozwiązania rozdzielają funkcje absorpcji strumienia danych od procesów analizy – pierwszy moduł zapisuje dane surowe, drugi przetwarza je jako graf uzupełniając o dodatkowe właściwości czy metadane. Integralną częścią architektury jest język zapytań umożliwiający komunikację pomiędzy użytkownikiem a silnikiem bazodanowym. Cypher wyróżnia się intuicyjną składnią opartą na graficznych symbolach odwzorowujących topologię grafu i wspiera tworzenie wzorców bezpośrednio korespondujących z modelem przechowywania. Z technicznego punktu widzenia parser języka dzieli zapytanie na ciąg operacji realizowanych przez planer zapytań, który dobiera strategię dostępu do danych optymalną dla konkretnego algorytmu lub struktury grafu.

Warto wspomnieć o roli komponentów odpowiedzialnych za interoperacyjność. MillenniumDB używa koncepcji grafy domenowe jako wspólnego mianownika pomiędzy różnymi typami modeli oraz standardów wymiany danych. Podejście to pozwala zachować spójność semantyczną przy integracji różnych źródeł i jednocześnie utrzymać efektywność działania całego systemu. Umożliwia to budowę aplikacji korzystających równocześnie z cech grafu właściwości (bogate atrybuty) oraz RDF (zgodność ze standardem SPARQL). Architektura systemu grafowego może też uwzględniać warstwę buforowania – przykładem jest OrientDB wyposażony w wielowarstwowy system buforów równoważący obciążenie między pamięcią główną a dyskiem. Takie rozwiązanie ogranicza liczbę kosztownych operacji wejścia/wyjścia potrzebnych przy dużej liczbie intensywnych zapytań. W zależności od implementacji może ono obejmować cache wyników zapytań, cache struktur indeksowych czy nawet cache części topologii grafowej.

Ostateczny projekt architektury powinien wyważyć kompromis pomiędzy elastycznością modelowania a kosztami operacyjnymi. System czysto natywny będzie szybki przy analizach relacji, ale mniej uniwersalny względem integracji z





aplikacjami spoza ekosystemu grafowego. Z kolei rozwiązanie hybrydowe zapewni wszechstronność kosztem przewidywalnej wydajności przy specyficznych operacjach eksploatacyjnych. Przy wyborze strategii wdrożeniowej warto uwzględnić również charakterystykę danych wejściowych – statyczne repozytoria dokumentacyjne mogą być obsługiwane inaczej niż dynamiczne logi rejestrujące interakcje tysięcy użytkowników per sekunda, oraz przewidywany sposób ich wykorzystania przez jednostki analityczne lub aplikacje produkcyjne.

3 Historia rozwoju baz danych

3.1 Ewolucja technologii bazodanowych

Dane hierarchiczne

Model danych hierarchicznych był jednym z pierwszych szeroko stosowanych rozwiązań w historii bazodanowej, opierającym się na drzewiastej strukturze powiązań pomiędzy rekordami. Dane zorganizowane były w formie węzłów rodzic–dziecko, gdzie każdy element posiadał dokładnie jeden nadrzędny rekord i mógł mieć wiele podrzędnych. Relacje w takim modelu miały stały kierunek od korzenia ku liściom, co zapewniało prostotę implementacji i intuicyjność przeszukiwania. W kontekście ówczesnych ograniczeń technicznych ten schemat pozwalał na szybki dostęp do danych przewidzianych w ustalonej ścieżce, lecz sprawiał trudności przy potrzebie realizowania zapytań wychodzących poza tę sekwencyjność – np. przy odnajdywaniu elementów powiązanych ze sobą tylko pośrednio, poprzez wielu różnych przodków.

Systemy hierarchiczne dobrze sprawdzały się w sytuacjach o silnej stabilizacji struktury – przykładem mogły być katalogi produktów firm, gdzie drzewo kategorii pozostawało niezmiennie przez lata. Relacyjne bazy danych, które pojawiły się później, zaczęły przejmować ich rolę dzięki większej elastyczności w formułowaniu zapytań i obsłudze skomplikowanych złączeń. W modelu hierarchicznym trudno było dekomponować dane do postaci umożliwiającej ich łatwe ponowne użycie w innych kontekstach bez powielania fragmentów drzewa – co prowadziło do redundancji informacji. Operacje takie jak aktualizacja wymagały często przechodzenia przez całe gałęzie, by zachować spójność. W praktyce strukturę hierarchiczną można formalnie zapisać jako graf ukierunkowany acykliczny (DAG), gdzie krawędzie reprezentują relacje rodzic–dziecko, a każdy wierzchołek poza korzeniem ma dokładnie jedną krawędź przychodzącą. Schemat ten jest łatwy do implementacji w silniku bazodanowym poprzez tablice wskaźników lub zagnieżdżone rekordy, jednak przy większej złożoności relacji jego możliwości stają się ograniczone – brak



wsparcia dla wielu rodziców czy powiązań między gałęziami oznacza konieczność stosowania duplikatów lub mechanizmów obejścia. Hierarchiczny zapis informacji wywodzi się także z naturalnych modeli organizacyjnych używanych przez instytucje – dokumentacja techniczna często miała układ rozdziałowy odpowiadający tej logice.

W zastosowaniach czysto transakcyjnych model był wydajny: operacje odczytu polegały na podążaniu predefiniowaną ścieżką od korzenia do rekordu docelowego. Jednak, kiedy pojawiało się zapotrzebowanie na analizy przekrojowe lub dynamiczne zmiany struktury, zalety stawały się mniej widoczne. Przykład użycia hierarchicznego podejścia można odnaleźć także we wczesnych implementacjach systemów katalogowania zasobów sieciowych czy organizacji struktur działów przedsiębiorstw. Każdy dział mógł zawierać sekcje podrzędne aż do poziomu indywidualnych stanowisk – reprezentowanych jako liście – bez potrzeby definiowania relacji bocznych. Ta prostota była zaletą tam, gdzie model miał pełnić funkcję odwzorowania statycznej struktury instytucji i nie wymagano nagłych zmian. W zestawieniu z nowoczesnymi grafowymi bazami danych konstrukcja hierarchiczna jawi się jako szczególny przypadek grafu o bardzo restrykcyjnej topologii. Brak cykli upraszcza algorytmy wyszukiwania i znacznie redukuje złożoność operacji przejścia, jednak kosztem możliwości ekspresji relacji bardziej złożonych niż czysta relacja nadrzędny–podrzędny.

Obecne systemy grafowe potrafią subsumować model hierarchiczny poprzez odpowiednią konfigurację krawędzi i etykiet, co umożliwia wykorzystanie tych samych algorytmów wyszukiwania w strukturach hybrydowych łączących cechy hierarchii i pełnego grafu właściwości. Krytycznym problemem była adaptacja takiego modelu do środowisk wymagających integracji z różnorodnymi źródłami danych. Strukturalna sztywność hierarchii utrudniała mapowanie modeli pochodzących ze świata RDF czy grafu właściwości, które dopuszczają wielokrotne typy relacji oraz tzw. edge-as-node. Oznaczało to, że rozbudowa aplikacji o nowe typy powiązań często wymagała przebudowy całego drzewa. Interesujące jest to, że wiele współczesnych platform – Neo4j czy GraphDB – mimo skupienia na grafach ogólnych może efektywnie symulować dane hierarchiczne dzięki dedykowanym zapytaniom wzorcowym i restrykcjom dotyczącym kierunku krawędzi.

Użytkownik może tworzyć subgrafy będące drzewami osadzonymi wewnątrz bogatszej sieci połączeń i wykonywać na nich szybkie operacje eksploracyjne bez rezygnowania z elastyczności pozostałych części modelu. Ostatecznie można powiedzieć, że dane hierarchiczne stanowią protoplastę części współczesnych koncepcji organizowania informacji, oferując prostotę oraz przewidywalność zachowania kosztem ograniczonej adaptacyjności i ekspresyjności relacji. Dzisiejsze bazy grafowe są w stanie integrować dawne zalety tego podejścia jednocześnie



eliminując jego kluczowe słabości dzięki bardziej ogólnym strukturom relacyjnym oraz wsparciu dla dynamicznych modeli danych.

Model encja-relacja

Model encja–relacja (E–R) należy do klasy koncepcyjnych metod modelowania danych, których celem jest uchwycenie istotnych bytów występujących w analizowanej domenie oraz powiązań pomiędzy nimi. Został zaproponowany przez Petera Chena w 1976 roku jako „Ujednolicony widok danych”. Istotą jego podejścia jest rozdzielenie opisu obiektów, enkapsulowanych jako encje, od opisu więzi między nimi, definiowanych jako relacje. Każda encja odpowiada konkretnej klasie obiektów o wspólnych cechach; może to być „Pracownik”, „Produkt” czy „Zamówienie”, natomiast relacje opisują ich wzajemne zależności, np. „pracuje w” lub „zawiera”. Formalny zapis modelu obejmuje typy encji, typy relacji oraz zestawy atrybutów charakteryzujących zarówno same obiekty, jak i powiązania. Tak skonstruowany opis jest niezależny od technologii przechowywania, pozwala odwzorować ten sam model zarówno na schemat relacyjny, jak i na strukturę grafową. W swojej klasycznej postaci diagram E–R przedstawia encje jako prostokąty, relacje jako romby oraz atrybuty jako elipsy przyłączone do odpowiednich elementów. Relacje mają nazwy semantycznie określające charakter powiązania i mogą posiadać kardynalność wskazującą liczby możliwych uczestnictw danego typu encji w relacji (np. 1:1, 1:N, N:M). W pierwotnym ujęciu możliwe było definiowanie jedynie pojedynczych nazwanych relacji pomiędzy dwiema encjami; model tym samym ograniczał możliwości ekspresji więzi o bogatej semantyce obecnych w złożonych domenach. Co istotne, E–R zakłada, że wszystkie powiązania są nieukierunkowane, kierunek wynika dopiero z interpretacji na poziomie implementacji.

Technologia bazodanowa lat osiemdziesiątych i dziewięćdziesiątych uczyniła z E–R podstawę projektowania schematów relacyjnych dzięki istnieniu dobrze określonych reguł translacji diagramów na struktury tabel. Od strony teoretycznej Codd wyprowadził równoważność między rachunkiem relacyjnym a algebrą relacyjną, co oznaczało, że każdą strukturę odwzorowaną z E–R można opisać formalizmem pierwszego rzędu i obsłużyć językiem SQL. Projektant mógł zatem pracować na modelu koncepcyjnym i mieć gwarancję jego pełnej reprezentowalności w systemach RDBMS. Kolejne etapy procesu obejmowały normalizację schematu zgodnie ze znanymi zasadami redukcji redundancji, etap ten był integralnym elementem pracy nad modelem w środowisku relacyjnym.

Interesującym aspektem ewolucji jest fakt, że struktura E–R może być systematycznie konwertowana także do schematu grafowego. Daje to projektantom możliwość zachowania dotychczasowych praktyk modelowania koncepcyjnego przy



jednoczesnym wyborze innej technologii przechowywania danych. Konwersja polega na potraktowaniu typów encji jako etykiet wierzchołków grafu właściwości oraz typów relacji jako etykiet krawędzi łączących odpowiednie węzły. Atrybuty stają się właściwościami przypisanymi albo bezpośrednio do węzłów (dla cech obiektu), albo do krawędzi (dla cech powiązania). Metoda ta wymaga świadomego wyboru nazw etykiet i kluczy dla właściwości tak, by zachować spójność semantyczną względem oryginalnego diagramu. Choć translacja jest teoretycznie prosta, praktyka ujawnia ograniczenia pierwotnego modelu E–R względem potrzeb dzisiejszych aplikacji grafowych. Diagram E–R dopuszcza tylko jedną nienazwaną kierunkowość relacji pomiędzy daną parą typów encji; brak możliwości tworzenia wielu polisemicznych więzi o różnych znaczeniach był akceptowalny w kontekście tabelarycznym, ale stanowi utrudnienie przy bezpośrednim odwzorowaniu na graf z licznymi semantycznymi krawędziami.

W świecie grafów właściwości jedna para obiektów może być połączona wieloma krawędziami o odrębnych typach i zestawach właściwości; aby odwzorować to z diagramu E–R należałoby rozbić pojedynczą logiczną więź na kilka typów relacji lub stosować mechanizmy krawędzi jako węzeł. Z historycznego punktu widzenia atrakcyjność modelu E–R wynikała z jego neutralności wobec implementacji: opisywał wyłącznie konceptualną strukturę danych bez narzucania sposobu jej zapisu w fizycznej bazie. To pozwalało różnym zespołom dyskutować wymagania biznesowe niezależnie od przyjętej technologii, szczególnie ważne było to przy integracjach systemowych między organizacjami stosującymi odmienne platformy bazodanowe. Diagram służył też jako narzędzie dokumentacyjne ułatwiające komunikację między działem analityki a programistami.

Współczesne grafowe bazy danych potrafią zaadaptować logikę E–R poprzez mapowanie jego komponentów na elementy graficzne: prostokąt odpowiada etykietie węzła, romb staje się etykietowaną krawędzią ukierunkowaną lub nieukierunkowaną zależnie od semantyki połączenia. Kardynalność zostaje przełożona na reguły integralności aplikowane podczas dodawania nowych krawędzi; można też wzbogacić model o metadane dotyczące dat utworzenia lub innych parametrów relacji. Różnice strukturalne powodują jednak zmianę perspektywy: tam, gdzie wcześniej istniała pojedyncza abstrakcyjna linia między encjami, grafowy zapis może obejmować wiele równoległych ścieżek reprezentujących rozmaite aspekty tej samej więzi. Przykład zastosowania konwersji można znaleźć w systemach zarządzania infrastrukturą IT: klasyczny model E–R opisuje zasoby takie jak „Serwer”, „Aplikacja” i „Użytkownik”, związane poprzez relacje „hostuje” czy „używa”. Przeniesienie tego modelu do bazy grafowej pozwala reprezentować dodatkowe połączenia wynikające z monitoringu wydajności lub zależności konfiguracji bez



przeprojektowywania całej struktury, wystarczy dodać nowe typy krawędzi łączące już istniejące węzły. Niektórzy projektanci zachowują hybrydowe podejście: tworzą diagramy E–R dla celów analizy wymagań biznesowych i planowania integralności danych zgodnie ze standardami SQL, ale implementują je od razu w formacie grafowym ze względu na przewagę tego modelu przy wykonywaniu zapytań dotyczących sieci połączeń czy hierarchii powiązań. Takie podejście łączy dyscyplinę projektową wykształconą przez dekady pracy z RDBMS z możliwościami technicznymi nowoczesnych konstrukcji najpierw graf (graph-first). Model E–R pozostaje więc istotnym punktem odniesienia zarówno dla historycznych analiz ewolucji modeli danych, jak i współczesnych praktyk integrujących różne paradygmaty, ze świadomością jego zalet (klarowność struktury koncepcyjnej) oraz ograniczeń (uboga ekspresja semantyczna więzi względem pełnych modeli grafowych) wynikających ze specyfiki czasów i technologii, w których powstał.

3.2 NoSQL i przejście do grafów

Na tle wcześniejszych modeli danych, które koncentrowały się na ścisłej strukturze hierarchicznej lub tabelarycznej, rozwój technologii NoSQL jawi się jako reakcja na rosnące wymagania w zakresie wolumenu, różnorodności i tempa napływu informacji. W poprzednich rozwiązaniach relacyjne systemy bazodanowe sprawnie obsługiwały spójne, znormalizowane dane, ale zaczynały tracić wydajność, gdy aplikacje wymagały częstych zmian schematu, pracy z danymi o luźnej strukturze czy obsługi bardzo dużych zbiorów powiązań. Rosnące zapotrzebowanie na obsługę nietypowych formatów oraz szybkie odpowiedzi przy wysokim obciążeniu doprowadziło do rozkwitu ekosystemu NoSQL obejmującego między innymi bazy dokumentowe, bazy klucz–wartość oraz modele kolumnowe. Każdy z tych typów unikał tradycyjnych operacji JOIN zastępując je mechanizmami bezpośredniego dostępu do rekordów, co poprawiało skalowalność horyzontalną kosztem możliwości reprezentowania relacji pierwszego rzędu.

Z perspektywy projektowej oznaczało to rezygnację z formalnego modelu encja–relacja opisanego w części wcześniejszej na rzecz form przechowywania umożliwiających łatwe dodawanie pól czy całych struktur bez konieczności migracji całej bazy. Jednak wiele zastosowań – szczególnie w analizach sieci społecznych czy systemach rekomendacyjnych – wymagało bardziej naturalnego sposobu odwzorowania połączeń i ich właściwości. W bazach dokumentowych relacja między obiektami była często redukowana do identyfikatorów umieszczonych wewnątrz dokumentu, bez natywnej obsługi połączenia jako odrębnego zasobu. To ograniczało



możliwość wykonywania dynamicznych zapytań łączących dane według wzorców topologicznych.

Grafowe bazy danych pojawiły się jako jeden ze specjalistycznych typów w rodzinie NoSQL, definiując swoją funkcjonalność wokół przetwarzania silnie powiązanych danych. Ich kluczową cechą stało się sąsiedztwo bez indeksowe – mechanizm pozwalający każdemu wierzchołkowi przechowywać bezpośrednie wskaźniki do swoich sąsiadów. Taka architektura upraszcza zapytania o skomplikowanych wzorcach relacji, które w środowiskach dokumentowych czy kolumnowych wymagałyby wielu osobnych operacji filtrowania i agregacji. W praktyce umożliwia to szybkie przechodzenie wielostopniowych ścieżek oraz analizę globalnych właściwości grafu. NoSQL dał ramy dla eksperymentowania z modelami przechowywania mniej restrykcyjnymi niż relacyjne RDBMS.

Grafowe systemy dopasowały się do tej dynamiki dodając jednocześnie ekspresyjny język zapytań dla struktury grafowej. Cypher stał się sztandarowym przykładem takiego podejścia – jego składnia inspirowana SQL pomagała użytkownikom znającym tradycyjne systemy przejść płynnie w stronę grafowego myślenia. Proste wyrażanie wzorców '(a)-[:REL]->(b)' odpowiada naturalnemu opisowi zależności obecnemu już w projektowaniu E–R, ale przenosi je na poziom operacyjny uwzględniający ukierunkowanie krawędzi i atrybuty relacji. Równocześnie kultura skalowania wypracowana w środowisku NoSQL odcisnęła piętno na projektach grafowych. Rozproszone magazyny musiały radzić sobie z problemem podziału grafu na fragmenty – tzw. sharding – unikając sytuacji, w której często używane krawędzie przekraczają granice fizycznych maszyn. Problem minimalnego punktu przecięcia oddaje istotę wyzwania: uzyskanie optymalnego rozdzielenia elementów przy minimalnej liczbie „przecinanych” relacji może być stanem chwilowym i zmieniać się wraz z ewolucją danych.

Systemy takie jak Neo4j rozwijały metody minimalizujące wpływ rozproszenia na czas odczytu poprzez replikacje fragmentów kluczowych lub inteligentne rozmieszczenie powiązanych węzłów. Nie można pominąć faktu, że pewna część funkcji grafowych była testowana najpierw w ramach narzędzi analitycznych działających poza bazami NoSQL – przykładem jest Google Pregel oraz jego wariant Giraph wykorzystywany w trybach OLAP do analizy dużych statycznych grafów. Rozwiązania te pracowały na klastrach i przetwarzały dane wsadowo; brak im było możliwości bieżącej modyfikacji danych czy wykonywania interaktywnych zapytań ad hoc. Integracja podejścia graph-first z doświadczeniem skalowalności NoSQL pozwoliła stworzyć narzędzia spójne dla obu trybów: analityki wsadowej oraz transakcyjnego przetwarzania zależności.



Ciekawy jest wpływ ruchu NoSQL na sposób kategoryzowania typów grafów stosowanych produkcyjnie. Stało się możliwe wdrażanie hybryd grafu właściwości i RDF triple store w jednym środowisku – część danych zachowywała wysoką szybkość przeszukiwania dzięki natywnej strukturze właściwościowej, a inna pozostawała kompatybilna ze standardami semantycznymi Internetu. MillenniumDB ilustruje tę tendencję poprzez model grafu domenowego łączący zalety RDF (interoperacyjność) i grafu właściwości (bogata semantyka atrybutów) przy jednoczesnym utrzymaniu zgodności operacyjnej między różnymi źródłami. Należy zauważyć, że migracja od klasycznych magazynów do grafowych rozwiązań często odbywała się etapowo – organizacje zachowywały dotychczasowe repozytoria dokumentowe lub kolumnowe jako pomocnicze źródła danych, podczas gdy krytyczne obszary analizy relacji przenoszono do dedykowanej bazy grafowej. Takie podejście redukuje ryzyko związane ze zmianą technologii i pozwala zespołom budować kompetencje grafowe równoległe z utrzymaniem dotychczasowych procesów.

Model NoSQL ułatwił też implementację mechanizmów odpornościowych potrzebnych przy wysokich wolumenach połączeń – architektury master-slave lub multi-master stosowane wcześniej przy magazynach kolumnowych znalazły adaptację przy klastrowaniu baz grafowych. Zachowanie spójności transakcji między partiami powiązań rozproszonych na wielu serwerach pozostało jednak wyzwaniem technicznym wymagającym kompromisów pomiędzy integralnością a dostępnością (zgodnie z zasadami teorii CAP). Z perspektywy dzisiejszych projektów można powiedzieć, że koncepcja NoSQL przygotowała grunt pod bardziej radykalny przewrót w sposobie myślenia o danych powiązanych, przejście do modeli graficznych wykorzystujących zalety swobodnego schematu i rozproszonego przetwarzania stało się naturalnym krokiem dla aplikacji operujących na gęsto splecionych sieciach informacji. Zasady skalowania i elastyczność struktury wyniesione ze środowisk dokumentowych czy kolumnowych zostały wzbogacone o natywne traktowanie relacji jako zasobu równorzędnego wobec samych obiektów, co czyni grafowe bazy danych pełnoprawnym elementem dużej rodziny technologii NoSQL ze swoimi własnymi specjalizacjami i scenariuszami użycia.

3.3 Wpływ rozwoju technologii na zastosowania grafów

Rozwój technologii bazodanowych, w tym postęp w architekturach natywnych i mechanizmach przetwarzania grafów, miał istotny wpływ na zakres oraz sposób ich zastosowań w różnych dziedzinach. Wraz z pojawieniem się rozwiązań typu graph-first, opartych na natywnym przechowywaniu i przetwarzaniu danych grafowych,



możliwe stało się znacznie szybsze wykonywanie operacji eksploracyjnych – każda krawędź przechowuje bezpośrednio odniesienia do sąsiadujących węzłów, co eliminuje konieczność korzystania z globalnych indeksów. Ta zmiana technologiczna przełożyła się na praktycznie natychmiastowe wyniki dla zapytań wymagających przechodzenia przez wiele poziomów relacji, co wcześniej w systemach nienatywnych wiązało się z poważnym wzrostem czasu odpowiedzi.

Rozszerzenie możliwości przechowywania metadanych zarówno na poziomie wierzchołków jak i krawędzi – typowe dla modelu grafu właściwości – umożliwiło budowę aplikacji, które nie tylko analizują strukturę połączeń, ale także wykorzystują kontekst opisowy relacji. Dzięki temu systemy rekomendacyjne mogą generować propozycje produktów uwzględniając nie tylko ścieżki połączeń między użytkownikiem a innymi obiektami, ale też parametry takie jak częstotliwość interakcji czy preferencje odnotowane w historii. Technologie grafowe znalazły zastosowanie w telekomunikacji i zarządzaniu sieciami komputerowymi. Mechanizmy takie jak korelacja zdarzeń pozwalają wiązać zdarzenia sieciowe o niskim poziomie abstrakcji z wyższego rzędu incydentami, a następnie oceniać ich wpływ na infrastrukturę poprzez analizę topologii grafu. Skrócenie czasu od wykrycia problemu do wdrożenia działań kompensujących jest efektem właśnie technologicznej optymalizacji w zakresie przechowywania i dostępu do powiązań – wcześniej operacje te mogły trwać godziny lub dni.

Z punktu widzenia analityki sieci społecznych czy systemów detekcji nadużyć kluczowe okazały się także algorytmy grafowe zoptymalizowane pod kątem dużych zbiorów danych. Mimo że skalowanie grafów rozproszonych jest nadal trudniejsze niż modeli kolumnowych czy dokumentowych, rozwój technik przetwarzania w czasie rzeczywistym umożliwił stosowanie bardziej skomplikowanych procedur np. wyszukiwania wzorców czy analiz przepływu bez opóźnień charakterystycznych dla trybów wsadowych. Platformy wyposażone w pamięci operacyjne oraz wielowątkowe procesory pozwalają dziś prowadzić analizy trójek RDF w czasie rzeczywistym. Ewolucja sprzętu również miała swoje odbicie w zastosowaniach – zastąpienie dysków talerzowych urządzeniami SSD czy pamięcią flash klasy enterprise pozwoliło zmniejszyć latencję dostępu do danych nawet dwudziestokrotnie. W kontekście baz grafowych oznacza to możliwość równie szybkiego odczytu fragmentów nie mieszczących się w pamięci głównej, co zwiększyło praktyczne pole wykorzystania dla scenariuszy wymagających analizy dużych historycznych zbiorów połączeń. Integracja modeli hybrydowych, takich jak grafy domenowe łączących RDF z grafami właściwości, otworzyła drogę do zastosowań wymagających interoperacyjności między różnymi standardami wymiany informacji. Tego rodzaju konstrukcja mechanizmów krawędzi jako węzeł czy zagnieżdżone węzły krawędziowe poszerza



modelowanie o relacje drugiego rzędu, np. opisujące same powiązania jako byty posiadające dodatkowe własności. Umożliwia to budowę bardziej zaawansowanych systemów zarządzania wiedzą w przedsiębiorstwach. Widoczny jest również wpływ języków zapytań takich jak Cypher na adopcję technologii grafowych w nowych obszarach. Znajoma semantyka podobna do SQL ułatwiła zespołom migrację z tradycyjnych środowisk bazodanowych.

Możliwość intuicyjnego formułowania wzorców struktury relacyjnej pozwala programistom szybko wdrażać modele dla takich zastosowań jak optymalizacja tras logistycznych czy monitorowanie stanu sieci sprzedaży. Postęp technologiczny umożliwił ponadto wykorzystanie algorytmów bardziej zaawansowanych niż proste zapytania typu k-hop. Dziś wsparcie dla pełnego odzyskiwania ścieżki czy regularnego zapytania o ścieżkę daje narzędzia potrzebne przy eksploracji zależności wieloetapowych – co ma zastosowanie choćby przy analizie quasi-hierarchicznych struktur organizacyjnych adaptowanych do modelu grafowego. Jedną ze zmian o wymiarze biznesowym jest to, że technologie grafowe przestały być domeną wyłącznie specjalistycznych zastosowań; rosnąca wydajność i elastyczność sprawiają, że trafiają one do codziennych procesów przetwarzania danych powiązanych w sektorach takich jak e-commerce czy administracja publiczna. Wdrożenie mechanizmów ACID przy natywnym zapisie zwiększyło atrakcyjność grafowych baz danych także tam, gdzie rygory spójności transakcyjnej są kluczowe.

Rozwój metod klastrowania i rozproszonego zarządzania zasobami zapewnił możliwość obsługi bardzo dużych wolumenów połączeń bez utraty dostępności systemu. Mechanizmy master-slave oraz konfiguracje multi-master znane wcześniej z innych magazynów NoSQL zostały zaadaptowane do wymogów spójności sieci grafowej przy równoczesnym zachowaniu odporności na awarie. Takie podejście rozszerza pole zastosowania o segmenty wymagające wysokiej niezawodności infrastruktury IT.

Warto zauważyć transformację sposobu myślenia o samym modelu danych – pojawienie się grafu domeny właściwości jako odmiany grafu właściwości pozwala zachować kompatybilność ze źródłami o różnym formacie atrybutów bez utraty wydajności. Zmiany te mają bezpośredni wpływ na zdolność systemu do integracji informacji pochodzących z wielu miejsc i prezentowania ich jako jednolitej struktury relacyjnej. Efekt synergii między rozwojem architektur natywnych, wzrostem mocy obliczeniowej sprzętu oraz udoskonaleniem języków zapytań jest wyraźny: bazy grafowe mogą dziś wykonywać kompleksowe analizy oparte zarówno na topologii relacji, jak i ich szczegółowym kontekście opisowym niemal natychmiast po otrzymaniu zapytania. To przełożyło się na nowe sposoby wykorzystania tej



technologii – od monitoringu cyberbezpieczeństwa poprzez śledzenie logiki interakcji klient–produkt aż po badanie semantycznie powiązanych repozytoriów wiedzy wspierających procesy decyzyjne organizacji.

4 Modelowanie danych w bazach grafowych

4.1 Reprezentacja danych w postaci grafu

Reprezentacja danych w postaci grafu opiera się na potraktowaniu wszystkich elementów systemu jako węzłów (wierzchołków) powiązanych między sobą krawędziami, które pełnią rolę odwzorowania relacji. W takim modelu każdy wierzchołek jest jednoznacznie identyfikowany i może być opisany przez zestaw właściwości w formie par klucz–wartość. Klucz definiuje nazwę cechy, a wartość jej konkretną treść – tekstową, liczbową czy też inną dopuszczalną przez magazyn danych. Ten sam mechanizm obowiązuje dla krawędzi, co pozwala zachować informacje kontekstowe na poziomie relacji, takie jak daty utworzenia połączenia, jego intensywność czy wagę istotną z punktu widzenia późniejszych analiz.

Podstawowy formalny zapis grafu to para zbiorów $G=(V,E)$, gdzie V stanowi zbiór wszystkich węzłów, a E zbiór wszystkich krawędzi. Każda krawędź jest relacją pomiędzy dwoma obiektami (v_out , v_in) uzupełnioną o kierunek i etykietę opisującą rodzaj zależności. W modelach typu graf właściwości relacje są zawsze skierowane – nawet jeśli semantyka sugeruje brak kierunkowości – ponieważ kierunek w strukturze pozwala precyzyjnie określić punkt startowy i docelowy zapytania. W przypadku potrzeby odwzorowania połączenia symetrycznego można ignorować ten kierunek na poziomie logiki zapytań. Etykiety są integralnym elementem reprezentacji danych. Wierzchołkom przypisuje się jedną lub więcej etykiet grupujących je według przeznaczenia lub kategorii. Przykład: etykieta User identyfikuje obiekty będące użytkownikami systemu, podczas gdy Product odnosi się do zasobów oferowanych przez firmę. Równolegle typ krawędzi zapisany jako etykieta relacji może mieć nazwę PURCHASED, LIKES czy CONNECTED_TO, co pozwala jednolicie interpretować semantykę grafu.

Systemy grafowe umożliwiają stosowanie wielu etykiet dla jednego elementu – zwiększa to elastyczność filtrowania i dostosowania modelu do różnych kontekstów analitycznych. Sposób odwzorowania bardziej złożonych relacji – takich jak powiązania wieloelementowe – wymaga rozszerzenia podstawowego schematu binarnego. Jedną z metod jest użycie konstrukcji krawędź jako węzeł: zamiast tworzenia pojedynczej krawędzi łączącej wszystkie uczestniczące byty, powołuje się nowy wierzchołek reprezentujący samą relację i łączy go odrębnymi krawędziami ze



wszystkimi uczestnikami. Takie podejście pojawia się m.in. w grafach domenowych oferujących kompatybilność z RDF poprzez przechowywanie wyraźnych identyfikatorów krawędzi, co pozwala dodawać własne atrybuty do konkretnej instancji połączenia bez utraty spójności semantycznej. Model RDF przedstawia dane jako trójki ("podmiot", "predykat", "obiekt") odpowiadające odpowiednio węzłowi źródłowemu, nazwie relacji oraz węzłowi docelowemu. Każdy fakt jest osobnym rekordem i budowa większych struktur wymaga interpretacji wielu trójek naraz. Zaletą RDF jest standaryzacja interfejsów dostępu takich jak SPARQL; wadą względem grafu właściwości może być niższa wydajność przy eksploracji głęboko powiązanych fragmentów grafu z uwagi na konieczność rekonstrukcji połączeń z kilku niezależnych elementów.

Z technicznego punktu widzenia reprezentacja danych obejmuje także mechanizm fizycznego powiązania rekordu wierzchołka z rekordem krawędzi. Przy natywnym przechowywaniu sąsiedztwo bez indeksów oznacza istnienie bezpośrednich wskaźników do sąsiadów, które eliminują potrzebę wyszukiwania ich poprzez globalne indeksy. W implementacjach nienatywnych struktura graficzna bywa odtwarzana z tabeli lub pliku dokumentowego – zwiększa to koszt operacyjny i wpływa na czas odpowiedzi przy zapytaniach obejmujących wiele skoków po grafie. Przedstawienie obiektów jako faktorów zawierających zarówno identyfikator, jak i zestaw właściwości sprzyja skalowalności modelu. Można łatwo dodawać nowe pola do istniejących elementów bez modyfikowania całego schematu bazy. Jest to szczególnie widoczne przy integracjach z różnorodnymi źródłami danych – model grafu domeny właściwości pozwala np. odwzorować właściwości tekstowe jako odrębne połączenia trójek ("id", "typAtrybutu", "wartość"), co zachowuje jednolity format metadanych nawet dla źródeł o niejednorodnych wariantach zapisu.

Reprezentacja grafowa dopuszcza też stosowanie specjalnych typów połączeń jak tzw. wirtualne krawędzie, które usługowo tworzone są podczas wykonywania zapytań na podstawie reguł agregujących lub porównujących wartości właściwości dwóch wierzchołków. Choć nie istnieją fizycznie w magazynie danych, wymuszają pewną dyscyplinę projektową: należy przewidzieć ich wpływ na wydajność oraz poprawność wyników dla bardziej rozbudowanych operacji analitycznych. Bezpośrednie oddanie realnych obiektów i ich więzi w strukturach grafowych sprawia, że każde zapytanie może operować jednocześnie na topologii sieci oraz na zawartych wewnątrz niej metadanych opisowych. Ta integralność modelu powoduje, że aplikacje takie jak systemy rekomendacyjne czy platformy detekcji nadużyć mogą korzystać np. ze wzorców ścieżek o określonych parametrach liczbowych zawartych w atrybutach krawędzi. Ważnym elementem jest tu możliwość mieszania typologii grafowej: część relacji może być ukierunkowana dla precyzyjnych scenariuszy (np.



przepływy transakcji), a część pozbawiona kierunku tam, gdzie semantyka zakłada równorzędność obu stron. Projektowanie reprezentacji wymaga świadomego wyboru pomiędzy wydajnością przechowywania a bogactwem semantycznym danego rozwiązania. Zbyt duża liczba etykiet lub nadmierna szczegółowość właściwości może wpłynąć na koszt pamięciowy i czas indeksowania; równocześnie ograniczenie informacji kontekstowych redukuje wartość analityczną struktury. Znalezienie balansu między tymi aspektami jest jednym z kluczowych zadań przy konstruowaniu modelu danych opartego na grafach dla konkretnych zastosowań biznesowych lub naukowych.

4.2 Projektowanie schematu grafu

Identyfikacja węzłów i relacji

Proces identyfikacji węzłów i relacji w projekcie schematu grafu wymaga określenia, które elementy domeny problemu będą reprezentowane jako jednostki pierwszego rzędu, a które jako połączenia między nimi. Z praktycznego punktu widzenia oznacza to konieczność rozróżnienia bytów o samodzielnym znaczeniu od zdarzeń czy zależności definiujących interakcje pomiędzy takimi bytami. Węzeł powinien odzwierciedlać istotny obiekt modelowanej rzeczywistości, użytkownika systemu, dokument w repozytorium, produkt oferowany przez firmę, konto bankowe albo transakcję traktowaną jako niezależną jednostkę w analizie. Każdy z węzłów posiada unikalny identyfikator oraz zestaw właściwości opisujących jego cechy.

Dobór tych właściwości powinien wynikać z wymagań analitycznych: dla klienta e-commerce mogą to być dane demograficzne i historia zakupów, dla urządzenia w sieci telekomunikacyjnej, parametry techniczne i lokalizacja. Relacje definiują powiązania pomiędzy parami lub większym zbiorem węzłów. W klasycznym modelu grafu właściwości są one kompatybilne strukturą z węzłami, mają własny identyfikator, etykietę typu relacji oraz opcjonalne właściwości kontekstowe. Etykieta relacji powinna jednoznacznie określać semantykę powiązania; może opisywać działania („PURCHASED”), stany („IS_MEMBER_OF”) lub zależności techniczne („CONNECTED_TO”). W przypadkach, gdzie związane węzły mogą mieć wiele typów powiązań jednocześnie, konieczne jest precyzyjne rozdzielenie relacji na osobne kategorie. System Neo4j np. dopuszcza definiowanie wielu krawędzi pomiędzy tą samą parą wierzchołków, każda z własnym typem i właściwościami. Takie podejście pozwala modelować sytuacje wieloaspektowe bez utraty przejrzystości struktury.



Moment identyfikacji węzła lub relacji często wiąże się z decyzją projektową dotyczącą granicy abstrakcji: czy zdarzenie (np. zakup) powinno być relacją pomiędzy klientem a produktem, czy raczej osobnym węzłem połączonym krawędziami do obu bytów. Koncepcja krawędzi jako węzeł rozpowszechniona m.in. w grafach domenowych przewiduje możliwość reprezentowania relacji jako pełnoprawnych węzłów. Pozwala to dodawać atrybuty bezpośrednio do „bytu-relacji” takie jak czas trwania czy status realizacji procesu, co bywa trudniejsze do odwzorowania przy standardowej krawędzi.

Odrębne zagadnienie stanowi kierunkowość relacji, techniczna specyfika grafu właściwości zakłada zapis każdej relacji jako skierowanej od v_{out} do v_{in} . Nawet jeśli semantycznie więź jest symetryczna (np. „friendOf”), jej implementacja zawiera wskazanie początkowego i końcowego wierzchołka; dla zachowania neutralności można tworzyć odwrotne relacje lub stosować logikę zapytań ignorującą kierunek. Kierunkowość determinuje zachowanie zapytań trawersujących graf, przy analizie przepływów finansowych kierunki muszą odpowiadać rzeczywistym transferom między rachunkami. Identyfikacja grup obiektów realizowana jest poprzez etykiety przypisane do węzłów: Person, Organization, Device, itd., co ułatwia filtrowanie danych według kategorii oraz implementację optymalizacji indeksowych ograniczonych do wybranej klasy obiektów. Warto przewidzieć spójny system nazewnictwa etykiet i typów relacji dla całej bazy, unikanie redundancji semantycznej redukuje ryzyko pojawienia się niespójnych wyników zapytań. Identyfikatory elementów powinny być globalnie unikalne wewnątrz bazy grafowej; Neo4j stosuje numeryczne ID przypisywane podczas tworzenia rekordu. Niektóre platformy umożliwiają stosowanie identyfikatorów zewnętrznych (UUID lub kody domenowe) zgodnych ze źródłem danych, co przyspiesza proces integracji.

W procesie projektowym istotne jest też rozpoznanie potencjalnych właściwości krawędzi: mogą one zawierać wagę dla algorytmów najkrótszej ścieżki, limit przepustowości dla problemów przepływu sieciowego czy parametry czasowe dla analizy sekwencji zdarzeń. Takie decyzje wpływają później na wybór algorytmów analizujących graf i ich efektywność. Na etapie identyfikacji warto uwzględnić ograniczenia kardynalności, ile instancji danej relacji może istnieć pomiędzy tymi samymi typami węzłów. Relacje takie jak owns mogą mieć charakter wielokrotny (klient może posiadać wiele kont), podczas gdy inne (uses, owes) występują najwyżej raz pomiędzy konkretnymi paroma bytami. Strategia modelowania powinna także przewidywać obecność specjalnych lub wyliczanych połączeń. Wirtualne krawędzie działające wg reguł generowanych na podstawie wartości atrybutów są przykładem takich konstrukcji; mimo braku fizycznego zapisu wymagają uwagi pod kątem planowania zapytań wykorzystujących je intensywnie. Kompletna identyfikacja



obejmuje mapowanie domenowych konceptów na elementy grafowe oraz określenie powiązań między nimi wraz z właściwościami opisowymi.

Przy złożonych źródłach danych – np. przy integracji trójek RDF – ważne jest ustalenie spójnego odwzorowania trójek na hierarchię etykiet i identyfikatorów grafu docelowego tak, aby możliwe było zarówno odtworzenie oryginalnych faktów, jak i wydajne ich przeszukiwanie po strukturze grafu właściwości. Wyważenie poziomu szczegółowości przy przypisywaniu ról elementom struktury decyduje nie tylko o poprawności modelu logicznego, lecz także o wskaźnikach wydajności operacyjnej całego systemu grafowego.

Określanie typów relacji

Określenie typów relacji w schemacie grafu jest etapem ściśle powiązaniem z wcześniejszą identyfikacją węzłów i krawędzi, lecz koncentruje się na doprecyzowaniu semantyki powiązań oraz ich roli w modelu. Typ relacji to swoista etykieta opisująca znaczenie i kontekst połączenia pomiędzy dwoma lub więcej węzłami, a jego wybór wpływa bezpośrednio na logikę zapytań oraz skuteczność analiz. Klasyczne modele grafu właściwości traktują typ relacji jako nazwę krawędzi – zapisywaną często wielkimi literami – odpowiadającą funkcji jaką pełni ona w strukturze. Przykładami mogą być PURCHASED, FRIEND_OF czy CONNECTED_TO. Dobór właściwej nazwy jest istotny nie tylko dla czytelności schematu, ale też dla spójności semantycznej aplikacji. Relacje można klasyfikować według kierunkowości: ukierunkowane od v_"out" do v_"in" lub nieukierunkowane, gdy znaczenie powiązania nie zależy od strony inicjatora.

W systemach takich jak Neo4j techniczna implementacja każdej krawędzi zakłada kierunek – nawet dla relacji semantycznie symetrycznych – z tego powodu warto rozważyć tworzenie par relacji odwrotnych lub przyjęcie konwencji ignorowania kierunku podczas wyszukiwania. Kierunkowość ma kluczowe znaczenie przy modelowaniu przepływów procesów, transakcji finansowych albo sekwencji zdarzeń. Podczas określania typów relacji należy także uwzględnić ich kardynalność. W praktyce oznacza to zdefiniowanie czy dana para typów węzłów może być połączona pojedynczą instancją relacji, wieloma instancjami równoległe (przypadek multigrafu) czy też żadną. Przy projektowaniu sieci komunikacyjnej etykieta CONNECTED_TO może łączyć dwa urządzenia wieloma kanałami fizycznymi – każde z nich reprezentowane osobną krawędzią o dodatkowym atrybucie capacity czy medium. W przeciwieństwie do tego relacja IS_MEMBER_OF między osobą a zespołem może występować tylko raz, co upraszcza logikę przetwarzania.

Określanie typów obejmuje również decyzję o tym, które relacje będą materializowane w magazynie danych, a które pozostaną logiczne lub wyliczane.



Tzw. krawędzie wirtualne pozwalają tworzyć połączenia "na żywo" podczas zapytania na podstawie reguł dotyczących właściwości węzłów (np. wspólny atrybut lokalizacji). Ponieważ takie krawędzie nie zajmują miejsca w pamięci stałej, redukują koszty przechowywania przy jednoczesnym zwiększeniu wymagań obliczeniowych podczas wykonywania zapytań. Przy intensywnym ich wykorzystywaniu należy przewidzieć mechanizmy buforowania wyników, aby zmniejszyć czas odpowiedzi. Semantyka typów relacji może być rozszerzona przez mechanizm krawędzi jako węzeł znany m.in. z grafów domenowych. Traktując samą relację jako niezależny byt można przypisać jej własne właściwości: datę utworzenia, status czy rodzaj interakcji. Przykład: zamiast łączyć użytkownika i produkt poprzez krawędź PURCHASED wyposażoną jedynie w atrybut date, stworzymy węzeł Purchase powiązany osobno z klientem i produktem oraz wzbogacony o zestaw metadanych takich jak metoda płatności czy punkt sprzedaży.

Określając typy należy brać pod uwagę interoperacyjność modelu ze źródłami trójek RDF. Każdy typ relacji powinien mieć odwzorowanie na predykat trójki ("podmiot", "predykat", "obiekt"), co ułatwia import/eksport danych między różnymi systemami. Podobnie przy hybrydzie grafu właściwości – RDF zachowanie spójności nazw predykatów jest krytyczne dla poprawnego działania równocześnie zapytań SPARQL i Cypher. Znaczenie typów przejawia się także przy optymalizacji indeksowej. Wydajne filtrowanie według typu krawędzi umożliwia szybkie ograniczenie zbioru analizowanych połączeń do tych istotnych dla danego scenariusza biznesowego. W praktyce stosuje się mechanizmy indeksowania na typ relacji lub na kombinację typu i wartości wybranych właściwości – np. wszystkie krawędzie typu CONNECTED_TO o przepustowości powyżej 100 Mbps. W systemach takich jak Neo4j zaleca się utrzymywanie jednolitego słownika typów, który eliminuje ryzyko niespójnego opisu więzi (np. użycie równolegle PURCHASED i BOUGHT).

Powtarzalność nazewnictwa ułatwia budowę zapytań wzorcowych skanujących fragmenty grafu pod kątem określonej semantyki. Na etapie projektowym warto więc ustanowić zasady formatowania etykiet (pisownia, separator słów) i ich dokumentowania. W przypadku danych dynamicznych istotna jest także decyzja o wersjonowaniu relacji. Typ może zostać rozszerzony o atrybut określający wersję lub stan historyczny; alternatywnie można tworzyć odrębne typy odpowiadające kolejnym wersjom kontaktu między obiektami. Takie rozwiązanie zwiększa liczbę typów w bazie, lecz upraszcza analizy ewolucji powiązań w czasie. Rozważając zastosowania praktyczne warto wspomnieć o sieciach społecznościowych modelujących różne poziomy interakcji – przykłady to FOLLOWS, LIKES, COMMENTED_ON – każdy z typów niesie odmienną informację o zachowaniu



użytkownika i pozwala na budowę wyspecjalizowanych algorytmów rekomendacyjnych wykorzystujących tylko wybrane więzi.

W infrastrukturze IT analogiczne role pełnią np. `DEPENDS_ON` dla zależności aplikacyjnych czy `HOSTS` dla przypisania zasobów sprzętowych do usług. Ostatecznie proces określania typów relacji powinien uwzględniać zarówno semantykę domenową jak i wymagania techniczne platformy grafowej. Balans pomiędzy szczegółowością katalogu typów a prostotą jego utrzymania jest warunkiem efektywnego działania bazy, nadmierne rozbudowanie listy rodzi ryzyko fragmentarycznych analiz trudnych do scalania, natomiast zbytne uproszczenie powoduje utratę informacji istotnych dla logicznych wzorców zapytań. Umiejętne dobranie zestawu typów stworzy fundament pod wydajne algorytmy eksploracyjne i precyzyjne mechanizmy filtrujące struktury grafowe zgodnie z potrzebami aplikacyjnymi lub badawczymi.

5 Języki zapytań w bazach grafowych

5.1 Przegląd języków zapytań

Języki zapytań w bazach grafowych pełnią rolę warstwy komunikacyjnej pomiędzy użytkownikiem a silnikiem bazy, umożliwiając formułowanie pytań o strukturę, właściwości i powiązania danych zapisanych w formacie grafowym. Klasyfikując je w kontekście modelu grafu właściwości czy RDF, można zauważyć różnice w filozofii ich projektowania oraz zakresie obsługiwanych operacji. Cypher jest jednym z najbardziej rozpoznawalnych języków zapytań dla grafów typu graf właściwości i został stworzony przez twórców Neo4j jako składnia intuicyjnie odzwierciedlająca topologię grafu poprzez symbolikę nawiasów okrągłych i strzałek.

Zapytanie `'(a:Person name:'Jim')-[:KNOWS]->(b)'` jest czytelnym wzorcem wyszukiwania – nawiasy oznaczają węzły wraz z etykietami i właściwościami, a strzałki kierunek relacji. W połączeniu z klauzulami `'RETURN'`, `'WHERE'`, `'CREATE'` czy `'CREATE UNIQUE'` daje to możliwość zarówno eksploracji istniejących powiązań, jak i tworzenia nowych. W praktyce Cypher jest językiem deklaratywnym – użytkownik określa, jakie dane chce otrzymać, bez konieczności opisywania algorytmu ich pozyskania. To podejście przypomina SQL, co ułatwia przejście osobom o doświadczeniu w relacyjnych bazach danych. Wśród języków reprezentujących środowisko RDF najważniejszy jest SPARQL (SPARQL Protocol and RDF Query Language). Operuje on na trójkach podmiot–predykat–obiekt, odwzorowując każdy fakt jako niezależny wpis o określonym predykatcie.



Składnia SPARQL pozwala łączyć wzorce trójek w celu odnalezienia pożądanego zbioru rezultatów, a dzięki wsparciu dla standardów W3C nadaje się do integracji z danymi sieci semantycznej. Istotną cechą jest możliwość mieszania danych uporządkowanych i nieustrukturyzowanych przy użyciu ontologii i taksonomii opisujących znaczenie relacji oraz pojęć. SPARQL pozwala również realizować mapowanie danych pomiędzy różnymi repozytoriami oraz wykrywać fakty implicytne na podstawie reguł inferencji semantycznej. Jest więc kluczowy w scenariuszach wymagających standaryzowanej wymiany informacji.

Gremlin reprezentuje natomiast rodzinę języków imperatywno-deklaratywnych opartych na paradygmacie przechodzenia. Zapytania Gremlin są sekwencją kroków przetwarzanych iteracyjnie od tzw. źródła przejścia – może to być obiekt g reprezentujący kontekst produkcyjny lub programistyczne środowisko robocze. Składnia Gremlin pozwala stosować złożone strategie przechodzenia po grafie zarówno z wykorzystaniem indeksów, jak i bez nich. Charakterystyczna jest konstrukcja łańcuchowa, gdzie kolejne kroki selekcji, filtrowania i agregacji wykonywane są na strumieniu odwiedzanych elementów. Dzięki temu Gremlin elastycznie łączy możliwości eksploracji punktowej z analizą całych ścieżek.

W bardziej akademickich kontekstach pojawiają się języki takie jak DGQL (Domain Graph Query Language), który adresuje specyficzne potrzeby pracy na tzw. Grafy domenowe. DGQL obsługuje zapytania obejmujące regularne zapytania o ścieżkę, pełne odzyskiwanie ścieżek czy wzorce nawigacyjne łączące warunki dotyczące zarówno typów wierzchołków, jak i krawędzi. Przykładowo możliwe jest wyszukanie artykułów napisanych przez autorów będących aktualnym personelem danej organizacji – wymaga to połączenia zmiennej krawędziowej z identyfikatorem typu węzła w trakcie jednego przebiegu analizy wzorca. Taka elastyczność wynika m.in. z mechanizmu dołączania pomiędzy krawędziami a węzłami.

Porównawcze analizy wskazują, że różne języki zapytań oferują odmienne zestawy funkcjonalności. Cypher zapewnia pełną obsługę podstawowych wzorców grafowych (BGP) oraz filtracji wyników według wartości atrybutów; SPARQL wyróżnia się kompatybilnością ze standardem RDF oraz wszechstronnością przy integracji modeli ontologicznych; Gremlin zaś preferowany jest tam, gdzie potrzebne jest proceduralne sterowanie ścieżką przechodzenia lub gdy wymagana jest spójność narracji między krokami przetwarzania; DGQL celuje w bogate wzorce relacyjne uwzględniające semantyczne aspekty modeli domeny. W kontekście interoperacyjności między tymi narzędziami istotną rolę odgrywają hybrydowe rozwiązania – niektóre bazy wspierają równoległe użycie Cypher i SPARQL czy Gremlin. Dzięki temu można korzystać z zalet każdej składni bez konieczności migracji całego modelu danych między różnymi środowiskami.



Deklaratywność Cypher lub SPARQL skraca czas implementacji prostych zapytań analitycznych, natomiast imperatywny styl Gremlin bywa niezbędny przy niestandardowych algorytmach przejścia wymagających dodatkowej logiki warunkowej. Nie sposób pominąć aspektu optymalizacji wykonania zapytań. Silniki takie jak Neo4j budują wewnętrzne plany wykonania dopasowane do struktury grafu – wybór odpowiedniej strategii zależy m.in. od liczby dostępnych indeksów dla etykiet i właściwości elementów użytych we wzorcu. Optymalizacja ta różni się między językami: SPARQL musi często rekonstruować powiązania na bazie wielu trójek, podczas gdy graf właściwości interpretuje całość relacji jako pojedynczą krawędź zawierającą wszystkie niezbędne informacje kontekstowe.

Zastosowanie konkretnego języka wiąże się więc ze świadomym kompromisem pomiędzy ekspresyjnością semantyczną a wydajnością przetwarzania wybranych typów zapytań. Tam gdzie encje są opisane bogatymi ontologiami – przewagę ma SPARQL dzięki integralności ze światem semantycznym; gdy priorytetem jest szybka analiza gęstej sieci powiązań gospodarczych lub społecznych – Cypher dzięki sąsiedztwu bez indeksowemu zapewni minimalny czas dostępu do kolejnych poziomów relacji; jeśli potrzebna jest logika proceduralna integrująca analizę z wykonaniem działań modyfikujących strukturę – wybór pada zwykle na Gremlin; przy modelach domenowych wymagających rozszerzonego opisu samych krawędzi oraz mechanizmów krawędź jako węzeł warto zwrócić się ku DGQL z jego wszechstronnymi wzorcami. Ewolucja tych języków pokazuje trend ku zapewnieniu zgodności semantycznej między środowiskami oraz coraz lepszemu wsparciu zarówno dla trybów OLTP (transakcyjnego), jak i OLAP (analitycznego). Postęp idzie także w stronę automatycznej translacji wybranych konstrukcji – tak aby np. wzorec Cyphera można było częściowo przekształcić do SPARQL lub odwrotnie przy zachowaniu intencji pytania. W efekcie współczesne bazy grafowe coraz częściej stają się poliglotyczne pod względem obsługiwanej składni, co sprzyja ich adaptacji w organizacjach korzystających dotąd z różnych sposobów opisu swoich zasobów informacyjnych.

5.2 Cypher

Podstawowa składnia

Podstawowa składnia języka Cypher jest zbudowana w sposób mający odzwierciedlać naturalny sposób myślenia o danych zapisanych w postaci grafu. Polega na tworzeniu wzorców, w których węzły przedstawione są w nawiasach okrągłych, a krawędzie jako strzałki ' $\rightarrow$ ' lub ' $\leftarrow$ ', zgodnie z kierunkiem relacji. Węzeł można oznaczyć etykietą, zapisaną po dwukropku – przykładowo '(a:Person)' – oraz



opcjonalnymi właściwościami ujętymi w klamry, jak 'name:'Jim". Strzałka pomiędzy węzłami wymaga określenia typu relacji przez wpisanie go w nawiasach kwadratowych, np. '[: KNOWS]'. Tak skonstruowany wzorzec jest rdzeniem większości zapytań Cyphera i może występować w różnych kontekstach. Klauzula 'MATCH' służy do dopasowania takiego wzorca do istniejących danych w grafie. Składnia '(a:Person name:'Jim') - [: KNOWS] ->(b)' oznacza wyszukanie znajomych osoby o imieniu Jim i przypisanie ich do identyfikatora 'b'. Dopasowanie odbywa się względem struktury grafu oraz warunków filtrowania właściwości. Oprócz 'MATCH' istotna jest klauzula 'WHERE', której rola polega na dodatkowej filtracji wyników na podstawie wyrażeń logicznych. Można ją zastosować zarówno dla właściwości węzłów, jak i krawędzi, a także do warunków strukturalnych – na przykład sprawdzania braku pewnych powiązań.

Przeniesienie kryteriów wyszukiwania z części wzorcowej do 'WHERE' nie zmienia semantyki wyniku, lecz może wpływać na czytelność oraz plan wykonania zapytania: '(a)- [: KNOWS] ->(b) WHERE a.name='Jim"' odpowiada formalnie dopasowaniu z właściwością umieszczoną bezpośrednio przy definicji węzła. Po uzyskaniu dopasowania należy wskazać elementy, które mają zostać zwrócone – służy do tego klauzula 'RETURN'. Podajemy tutaj identyfikatory węzłów, krawędzi lub ich właściwości: 'RETURN b.name' zwróci tylko imiona osób powiązanych z Jimem znalezionych pod identyfikatorem 'b'. Możliwe jest zwracanie całych ścieżek ('RETURN path') lub zestawów wielu kolumn z różnych elementów wzorca. W wersjach wspierających agregację można tu stosować funkcje takie jak 'COUNT(b)' czy 'COLLECT(b.name)', co zwiększa zakres przetwarzania bez konieczności opuszczania kontekstu Cyphera.

Podstawowa składnia obejmuje też klauzule umożliwiające modyfikację danych. Klauzula 'CREATE' tworzy nowe węzły lub relacje zgodnie ze wskazanym wzorcem: 'CREATE (a:Person name:'Ian') - [: KNOWS] -> (b:Person name:'Emil')' utworzy oba obiekty oraz ich połączenie. Alternatywnie, 'CREATE UNIQUE' zapewni unikalność wzorca poprzez stworzenie połączenia tylko wtedy, gdy nie istnieje identyczny układ elementów i właściwości. Klauzula 'MERGE' łączy działanie wyszukiujące i tworzące – jeśli podany wzorzec istnieje, zostaje użyty; jeśli nie, zostaje stworzony nowy zestaw elementów spełniający warunki. Jest to mechanizm przydatny przy aktualizacjach wymagających zachowania spójności modelu. Modyfikacja właściwości to domena klauzuli 'SET', która pozwala dodać nową parę klucz–wartość lub zastąpić istniejącą: 'SET a.age = 34'.

Usuwanie danych realizowane jest poprzez klauzulę 'DELETE', mogącą usuwać pojedyncze elementy lub całe ścieżki zależnie od podanego kontekstu. Jeżeli chcemy usunąć zarówno węzeł, jak i jego połączenia, stosujemy dodatkowo słowo



kluczowe 'DETACH DELETE'. Składnia przewiduje też mechanizmy przetwarzania zbiorów wyników i łączenia wielu zapytań. Klauzula 'WITH' przekazuje wyniki jednego bloku zapytania jako wejście dla kolejnego – przypomina to potok (pipeline) znany z systemu Unix. Pozwala również definiować aliasy dla wartości obliczeniowych przed przejściem do kolejnej części przetwarzania. Łączenie wyników wielu zapytań wykonuje się za pomocą 'UNION', które scala dane o jednakowej strukturze kolumnowej.

W podstawowym zakresie składni trzeba także wspomnieć o dawnym słowie kluczowym 'START'. Choć zostało ono uznane za przestarzałe na rzecz osiągania punktów początkowych poprzez dopasowanie w klauzulach MATCH, przy wcześniejszych wersjach Neo4j można było wskazywać konkretny wierzchołek czy relację jako punkt początkowy według jej ID lub indeksu. Interesującym aspektem codziennych zastosowań podstawowej składni jest możliwość zwracania całych ścieżek zamiast pojedynczych elementów. Dzięki temu użytkownik może badać topologię powiązań w ich pełnym kontekście bez utraty informacji o kierunku czy typach relacji pomiędzy kolejnymi segmentami drogi.

Pełne dopasowanie mechaniki Cyphera wymaga też świadomości domyślnego działania tzw. leniwego wiązania – identyfikatory przypisywane są do kolejnych dopasowanych elementów stopniowo podczas iterowania wyników. Pozwala to redukować koszty pamięciowe przy dużych odpowiedziach i umożliwia strumieniowe przetwarzanie danych bez konieczności buforowania całości. Jak pokazuje praktyka korzystania z bibliotek typu MillenniumDB czy narzędzi typu RDF–hybryda własności 5.1, prostota tych konstrukcji sprawia, że można je łatwo rozbudowywać o bardziej zaawansowane formy filtrowania czy agregacji. Jednak fundament – nawiasy okrągłe dla węzłów, strzałki dla relacji i przejrzyste klauzule kontrolne ('MATCH', 'WHERE', 'RETURN') – pozostaje punktem wyjścia każdego zapytania analizującego strukturę grafową.

Zaawansowane konstrukcje

Rozszerzone możliwości języka Cypher pozwalają na budowę zapytań o bardziej złożonej logice niż proste wzorce dopasowania omówione wcześniej. W tej warstwie pojawiają się mechanizmy umożliwiające kontrolę przepływu danych w zapytaniu, zagnieżdżanie filtrów, operacje na ścieżkach oraz korzystanie z wyrażeń agregujących i podzapytaniach. Jednym z kluczowych elementów jest użycie klauzuli 'WITH' w kontekście zaawansowanym – może ona nie tylko przekazywać wynik jednego etapu przetwarzania do kolejnego, ale również modyfikować zestaw danych poprzez opcje filtrowania lub sortowania przed dalszą analizą. Przykładowo, można najpierw znaleźć wszystkich znajomych o określonej cenie, posortować ich według



atrybutu score, a następnie przekazać tylko pierwszych dziesięciu do kolejnej części zapytania.

Zaawansowane konstrukcje obejmują również wzorce krawędziowe o zmiennej długości. Użycie składni typu `'-[:REL*2...5]->'` pozwala dopasować ścieżki zawierające od dwóch do pięciu krawędzi określonego typu pomiędzy węzłem początkowym i końcowym. Mechanizm ten jest istotny przy poszukiwaniu powiązań o określonej głębokości w sieciach społecznych czy systemach rekomendacyjnych. W połączeniu z aliasowaniem ścieżek (`'MATCH p = (a)-[:REL*1..3]->(b)'`) można później wykorzystać funkcje operujące na nich, np. `length(p)`, co daje dodatkową kontrolę nad wynikami.

Ciekawym aspektem są funkcje agregujące stosowane w kontekście filtrowania wyników. Cypher oferuje takie operatory jak `COUNT()`, `COLLECT()`, `AVG()` czy `MAX()`, które mogą być używane zarówno w sekcji `'RETURN'`, jak i `'WITH'`. To umożliwia grupowanie rezultatów według wartości właściwości lub typów relacji. W praktyce można np. policzyć liczbę połączeń określonego rodzaju dla danego węzła i wyświetlić tylko te obiekty, które spełniają warunek minimalnej liczby takich relacji. Operacje warunkowe i predykaty logiczne stanowią kolejną kategorię konstrukcji rozszerzających.

Oprócz standardowych operatorów porównania Cypher obsługuje również predykaty istnienia (`'EXISTS()'`), które pozwalają sprawdzić obecność pewnych właściwości lub relacji w dopasowanym fragmencie grafu. Można np. znaleźć wszystkie osoby posiadające relację `WORKS_FOR` i jednocześnie pole `position` ustawione na „manager”. W pracy z dużymi zbiorami danych istotne staje się też filtrowanie poprzez wzorce negatywne, konstrukcja `'WHERE NOT (...)'` daje możliwość eliminacji elementów posiadających określone powiązania lub cechy, bez konieczności skomplikowanego przekształcania zapytania. Pozwala to efektywnie oczyszczać wyniki ze zbędnych elementów jeszcze przed etapem agregacji czy sortowania. Interesujące są także zastosowania mechanizmu rozumienia wzorców, który umożliwia tworzenie list wyników bezpośrednio wewnątrz pojedynczego wzorca dopasowania.

Składnia `'[x IN nodes(path) WHERE x.prop > 10 | x.name]'` generuje listę nazw tych węzłów ścieżki, które spełniają podany warunek. Takie podejście skraca czas potrzebny na obróbkę wyników po stronie aplikacji – logika wyboru elementów zostaje zamknięta w jednym wyrażeniu języka zapytań. Kolejnym krokiem w stronę zaawansowanych scenariuszy jest tworzenie podzapytań i praca na kontekstach lokalnych wynikających z częściowego dopasowania grafu. Podzapytania uruchamiane przez `'CALL ...'` mogą zwracać dane do głównego zapytania lub



realizować niezależne procesy analityczne. Pozwala to m.in. łączyć różne typy wyszukiwań – np. najpierw pobrać podgraf spełniający określony wzorec domenowy, a potem wykorzystać jego elementy jako wejście do algorytmu wyszukiwania najkrótszej ścieżki. Zaawansowana składnia wspiera również implementację koncepcji krawędź jako węzeł opisywanej dla grafów domenowych, krawędź może zostać potraktowana jako węzeł i uczestniczyć we wzorcach tak samo jak inne jednostki modelu. W Cypher realizuje się to poprzez sekwencję dopasowań uwzględniających identyfikator krawędzi oraz jej właściwości – daje to możliwość analizy meta-relacji pomiędzy połączeniami (np. porównywanie czasu trwania różnych typów interakcji).

Użytkownicy wymagający kontroli nad wydajnością mogą skorzystać z indeksowania per etykieta i właściwość, chociaż jest to aspekt konfiguracyjny bazy, wpływa na plan wykonania skomplikowanych zapytań poprzez redukcję kosztu wyszukiwania punktu początkowego dla przechodzenia. W połączeniu ze świadomym ustaleniem punktów wejścia ('MATCH (start:Label key:'value')') można znacznie skrócić czas reakcji przy pytaniach obejmujących duże fragmenty grafu. Nie bez znaczenia pozostaje możliwość mieszania Cyphera z innymi językami zapytań wspieranymi przez bazę, przykładowo MillenniumDB umożliwia współpracę składni typowej dla języka zapytań dotyczących grafów domenowych wraz ze specyficznymi cechami Cyphera. Takie hybrydowe podejście otwiera drogę do formułowania bardzo niestandardowych konstrukcji uwzględniających zarówno semantykę RDF/ SPARQL jak i ekspresję grafu właściwości.

Podsumowując techniczny wymiar zaawansowanych konstrukcji należy wskazać na rolę elastycznego łączenia wielu klauzul ('MATCH', 'WHERE', 'WITH', 'CALL') oraz dynamicznych wzorców o zmiennej długości krawędzi jako fundamentu budowy bardziej rozbudowanych zapytań. Ich umiejętne użycie daje możliwość tworzenia precyzyjnych narzędzi eksploracyjnych działających bezpośrednio na strukturze grafowej, co ma szczególne znaczenie przy badaniu gęsto powiązanych sieci informacyjnych czy modelowaniu procesów biznesowych o wieloetapowej logice.

5.3 Gremlin

Podstawowa składnia

Składnia języka Gremlin opiera się na podejściu do zapytań jako sekwencji kroków przetwarzania strumienia danych, gdzie wynik każdego kroku staje się wejściem dla kolejnego. Fundamentem jest koncepcja przechodzenia – operacji przechodzenia po grafie od węzła początkowego, przez krawędzie i kolejne węzły, aż do osiągnięcia



warunków końcowych. Każdy krok może być typu map (przekształcenie obiektów), filtr (usunięcie wybranych elementów) lub efekt uboczny (działanie uboczne jak np. zliczanie). Logika ta jest konsekwentna: w praktyce oznacza to budowanie łańcucha wywołań metod, które wyglądają jak komendy funkcyjnego języka skryptowego. Punkt startowy zazwyczaj określa zbiór wierzchołków lub krawędzi, przez które przechodzenie będzie przebiegało.

W kontekście systemów zgodnych z TinkerPop często przyjmuje to postać `'g.V()'` lub `'g.E()'`, gdzie `'g'` jest kontekstem grafu. Wywołanie `'g.V()'` wybiera wszystkie dostępne wierzchołki, natomiast przy podaniu parametru – identyfikatora albo właściwości – ogranicza je do tych spełniających warunek. Przykład prostej kwerendy "friend-of-a-friend" ilustruje typowy zapis Gremlina: `'v.out('friend').out('friend')`. Tu `"v"` reprezentuje wyjściowy punkt odniesienia, metoda `'out('friend')` wybiera sąsiadów po wychodzącej relacji o etykiecie "friend", a kolejne `'out('friend')` drąży dalej po tym samym typie powiązania. W ten sposób naturalnie tworzy się zapytania eksploracyjne, które w SQL wymagałyby złożonych złączeń.

Istotnym elementem składni są predykaty filtrujące – stosowane poprzez konstrukcję `'has(key,value)'` czy `'hasLabel(label)'`. Pozwalają one zawęzić zbiór bieżących elementów do tych posiadających określone właściwości lub należących do wskazanej kategorii etykietowej. Można je lokować zarówno na początku przechodzenia (co skraca ścieżkę przeszukiwania), jak i w trakcie późniejszych kroków dokładniej selekcionując dane na podstawie kryteriów domenowych. Warianty `'in ()'` oraz `'both()'` umożliwiają poruszanie się po relacjach odpowiednio przychodzących i dwukierunkowych – ich odpowiednie użycie jest kluczowe dla poprawnej interpretacji semantyki powiązań. Gremlin oferuje różnorodne mechanizmy sterowania długością przechodzenia. Jednym z nich jest użycie kroku `'repeat(...)'` wraz z `'times(n)'`, co pozwala wykonać sekwencję przejść określoną liczbę razy. Alternatywnie można zastosować `'until(condition)'` by powtarzać kroki aż do spełnienia warunku zakończenia – np. osiągnięcia węzła o danej właściwości. Takie podejście daje większą kontrolę niż prosta notacja Cyphera dla zmiennej liczby przeskoków, pozwalając implementować logikę bliską proceduralnemu opisowi algorytmu.

Budowa ścieżek i ich analiza odbywa się poprzez metody takie jak `'path()'`, która zachowuje pełny przebieg przechodzenia łącznie z odwiedzonymi elementami i krawędziami. Uzyskany obiekt ścieżki można następnie poddać dalszej obróbce przy użyciu standardowych kroków mapujących (`'map(...)'`) lub funkcji dostępu (`'values(...)'`, `'labels()'`). Dla analiz wymagających agregacji danych dostępne są kroki `'group()'`, `'groupCount()'`, które pozwalają formować słowniki wartości według wskazanego klucza lub liczyć wystąpienia poszczególnych cech. Koncepcja efekt



uboczny jest w Gremlinie istotna nie tylko dla generowania statystyk, ale także dla aktualizacji struktury grafu – kroki takie jak `' . addV(label)'` tworzą nowe wierzchołki, a `' .addE(label)'` dodają krawędzie między bieżącymi elementami a wskazanymi docelowo. Parametry tych metod zgadzają się z charakterem grafu właściwości: zarówno wierzchołkom jak i krawędziom można przypisywać właściwości przez kolejne wywołania `' . property(key,value)'` tuż po ich utworzeniu. Składnia Gremlina przewiduje także integrację z innymi paradygmatami programistycznymi – możliwe jest osadzanie zapytań wewnątrz kodu pisanego m.in. w Javie, Groovy czy Pythonie. To wyróżnia go spośród bardziej samodzielnych języków zapytań takich jak Cypher, dzięki temu logika aplikacyjna może bezpośrednio kontrolować przejście przepływu poprzez instrukcje warunkowe czy pętle implementowane w języku gospodarzu. Szczegółowość tej integracji sprawia, że translacja koncepcji SQL na model grafowy bywa nieintuicyjna, nie istnieje proste odwzorowanie JOIN na sekwencję kroków Gremlina.

W strukturze kroków często stosuje się równoległe przestrzenie nazw zwracanych elementów przez aliasowanie ich za pomocą metody `' .as('alias')`. Alias służy później do ponownego odwoływania się do tego samego punktu przechodzenia za pośrednictwem kroku `' . select('alias')`. Pozwala to „wracać” do wcześniejszego kontekstu bez konieczności restartowania całego procesu przeszukiwania od korzenia grafu. Ta możliwość szczególnie wspiera konstrukcje wymagające zestawienia cech kilku różnych segmentów ścieżki, np. porównania wartości właściwości dwóch niesąsiadujących bezpośrednio węzłów. Gremlin wspiera operacje analogiczne do projektowania widoków czasowych nad grafem dzięki metodom takim jak `' . subgraph('storageName')`, które zapisują podzbiór aktualnie odwiedzanych elementów jako osobny graf tymczasowy. Daje to narzędzia do segmentowania analizy dużych zbiorów połączeń, możliwe jest wykonanie intensywnej eksploracji lokalnej topologii bez ingerencji w cały globalny model danych.

W praktycznej pracy z Gremlinem warto pamiętać o wpływie rozmieszczenia filtrów na wydajność, umieszczenie kroku filtrującego wcześniej ogranicza liczbę elementów przekazywanych dalej, redukując koszt kolejnych etapów przechodzenia. Z drugiej strony filtry pośrednie mają znaczenie semantyczne, mogą być stosowane dopiero po uzyskaniu pewnych zależności między elementami budującymi ścieżkę. Całość podstawowej składni Gremlina sprowadza się więc do świadomego komponowania kroków przemieszczających strumień przez strukturę grafową przy jednoczesnym modyfikowaniu jej lub analizowaniu zgodnie ze specyfiką problemu badawczego. Jego elastyczność wynika między innymi z możliwości stosowania zarówno deklaratywnych wzorców bazujących na prostych selekcjach etykiet i właściwości, jak



i imperatywnych opisów algorytmu przejścia sterowanych przez konstruktory warunkowe oraz iteracje powtarzalne.

Wzorce zapytań

Wzorce zapytań w Gremlinie stanowią jedno z centralnych narzędzi umożliwiających efektywne przeszukiwanie i analizowanie danych zapisanych w strukturze grafowej. Opierają się one na sekwencjach kroków przechodzenia, które opisują sposób przechodzenia pomiędzy węzłami i krawędziami, uwzględniając warunki filtrowania oraz operacje transformacji strumienia wyników.

W praktyce wzorzec zapytania jest definicją ścieżki lub zestawu ścieżek w grafie o określonej topologii oraz właściwościach elementów. Kluczowym aspektem jest tu możliwość tworzenia wzorców wieloetapowych, w których operacje 'out()', 'in()' lub 'both()' budują złożone ciągi przejść pomiędzy etykietowanymi relacjami, a kroki filtrujące '.has()' czy '.where()' redukują przestrzeń wyszukiwania do tych fragmentów grafu, które mają znaczenie dla analizy. W odróżnieniu od deklaratywnej składni Cyphera, Gremlin wymaga świadomego projektowania kolejności kroków przechodzenia w taki sposób, aby wzorzec był zarówno semantycznie poprawny, jak i zoptymalizowany pod kątem wydajności. Przykładem może być wzorzec wyszukujący użytkowników połączonych bezpośrednio lub pośrednio poprzez relacje typu friend do konkretnego punktu startowego:

```
g.V().hasLabel('Person').has('name','Alice').repeat(out('friend')).times(2).path()
```

Tutaj sekwencja 'repeat(...).times(2)' definiuje długość dopasowywanej ścieżki w sposób iteracyjny; pozyskana metoda '.path()' umożliwia zachowanie pełnej historii przebytej trasy wraz z elementami spełniającymi warunki kwalifikacyjne. Wzorce mogą również stosować aliasowanie za pomocą '.as('alias')' w celu późniejszego odwoływania się do określonych punktów ścieżki, co bywa niezbędne przy konstrukcjach kontekstowych. Na przykład:

```
g.V().hasLabel('Order').as('o').out('contains').hasLabel('Product').select('o')
```

taki zapis pozwala zestawić właściwości węzła zamówienia z jego produktami nawet po przejściu kilku etapów przechodzenia. Alias działa jak uchwycony kontekst, można go pobrać z dowolnego miejsca późniejszej sekwencji kroków. Istotnym wariantem wzorców zapytań są te o zmiennej długości dopasowywane przez mechanizmy warunkowe '.until()' i '.repeat()', zamiast ustalonych długości skoku. Takie konstrukcje pozwalają na dynamiczne przeszukiwanie grafu aż do spełnienia wybranego kryterium końcowego:

```
g.V().hasLabel('Person').repeat(out()).until(has('status','inactive')).path()
```



Tutaj przechodzenie zatrzymuje się dopiero przy napotkaniu elementu spełniającego warunek dotyczący właściwości status. Dzięki temu uzyskujemy modele eksploracji odpowiadające realnym procesom biznesowym czy decyzyjnym. Gremlin obsługuje także podwójne wzorce filtrujące, etapowe zawężanie następuje w kilku miejscach traversala. W pierwszym kroku stosujemy szeroki filtr etykietowy lub własnościowy, a następnie wewnątrz wzorca ('where(...)') dodatkowe ograniczenia semantyczne lub topologiczne. Takie podejście jest szczególnie ważne przy analizach wymagających selekcji na podstawie zależności pomiędzy dwoma segmentami ścieżki, np. znajdowanie par urządzeń powiązanych krawędzią `CONNECTED_TO`, które jednocześnie znajdują się w tej samej lokalizacji geograficznej. Ważnym elementem projektowania wzorców jest także umiejętność wykorzystywania agregacji jako części konstrukcji przechodzenia. Gremlin umożliwia grupowanie wyników za pomocą `'group()'` i filtrowanie grup według kryteriów iteracyjnych, takie wzorce pozwalają opisać nie tylko pojedyncze ścieżki, ale też statyczne lub dynamiczne kolekcje powiązań według kontekstu biznesowego. Przykład:

```
g.V().hasLabel('Customer').group().by('country').by(outE('purchased')).count())
```

tworzy słownik liczby zakupionych produktów per kraj klienta.

Interesującą kategorią są wzorce oparte na subgrafach, umożliwiające utworzenie lokalnego modelu fragmentu grafu zgodnie z regułami przejścia zawartymi we wzorcu głównym. Konstrukcja `'subgraph('sg')'` zapisuje bieżący kontekst jako oddzielny obiekt do dalszej analizy:

```
g.V().hasLabel('Dept').outE().subgraph('sg')
```

W efekcie możliwe jest wykonywanie intensywnych obliczeń na ograniczonej części danych bez ingerencji w cały graf.

Niektóre przypadki wymagają wzorców korzystających ze zmiennych kontrolnych dostarczanych przez metodę `'select()'` i wielu aliasów naraz, takie konstrukcje wpływają na zdolność analizowania relacji nienależących do tego samego kontinuum przechodzenia, a jednak istotnych z perspektywy wynikowej odpowiedzi. To podejście sprawdza się w scenariuszach porównawczych czy korelacyjnych, gdzie zestawia się właściwości elementów odległych o kilka kroków lub należących do różnych gałęzi rozszerzonego przebiegu.

Projektując wzorce, warto brać pod uwagę koszt poszczególnych kroków i rozmieszczenie filtrów – ich wcześniejsze zastosowanie znacząco zmniejsza liczbę elementów badanych później. Świadome kompozycje kroków przechodzenia decydują o tym, czy otrzymany wynik będzie użyteczny i wydajny operacyjnie przy dużych zbiorach powiązań. Podsumowując techniczną stronę tematu, wzorce



zapytań w Gremlin to strategiczne kombinacje kroków o określonej kolejności i strukturze logicznej, które definiują proces eksploracji grafu oraz przekształcania jego danych zgodnie z wymaganiami analizy. Poprawne ich skonstruowanie wymaga nie tylko znajomości dostępnych metod języka, ale też umiejętności przewidywania skutków operacyjnych dla całego procesu przetwarzania dużych sieci powiązań.

6 Wydajność i skalowalność baz grafowych

6.1 Optymalizacja zapytań

Optymalizacja zapytań w bazach grafowych jest procesem, który wymaga zarówno zrozumienia wewnętrznej architektury systemu, jak i analizy topologii danych oraz sposobów ich przechowywania. Kluczowym elementem przyspieszającym wykonanie zapytań w środowiskach natywnych jest mechanizmu sąsiedztwa bez indeksowego, polegający na tym, że każdy wierzchołek posiada bezpośrednie wskaźniki do swoich sąsiadów. Dzięki temu eliminowana jest potrzeba globalnych wyszukiwań poprzez indeksy, co skraca czas przechodzenia przez wielopoziomowe relacje. W praktyce oznacza to, że zoptymalizowane zapytanie powinno startować od punktów grafu o możliwie ograniczonej liczbie połączeń wychodzących i stosować filtry już na początku ścieżki przeszukiwania.

Wdrożenie filtrowania na wczesnym etapie przechodzenia redukuje liczbę elementów poddawanych dalszym operacjom i zmniejsza obciążenie pamięci masowej. Istotnym aspektem optymalizacji jest świadome korzystanie z indeksów dla etykiet i właściwości – choć grafy natywne minimalizują ich rolę podczas przechodzenia, dobrze zaprojektowany indeks może znacząco skrócić czas wyszukiwania punktu startowego zapytania. W przypadku Neo4j możliwe jest utworzenie indeksu powiązanego z identyfikatorem zewnętrznym (np. kodem Q510 z Wikidata) lub konkretną właściwością domenową. Taka konstrukcja pozwala rozpocząć przechodzenie od ściśle określonego zestawu węzłów bez potrzeby pełnego skanu grafu. Jednak nadmiar indeksów – szczególnie dla rzadko używanych właściwości – może generować koszty związane z utrzymaniem spójności podczas operacji CREATE czy MERGE. Planowanie zapytań przez silnik bazy opiera się o koncepcję planowania zapytań, który wybiera strategię dostępu do danych optymalną dla danego wzorca.

Na wydajność wpływa zarówno struktura dopasowania ('MATCH' w Cypher), jak i kolejność łączenia jego segmentów przy użyciu 'WITH' czy 'OPTIONAL MATCH'. Przykład: najpierw zawężenie zbioru do małego fragmentu grafu na podstawie etykiety i właściwości ('MATCH (a:Person name:'Alice')'), następnie eksploracja



relacji wychodzących ('MATCH (a)-[:FRIEND]->(b)'), a dopiero później agregacja wyników. Odwrócenie kolejności może nie tylko wydłużyć czas wykonania, ale też spowodować przeciążenie pamięci ze względu na większy strumień danych przekazywany między etapami. Optymalizacja obejmuje również redukcję kosztownych joinów między krawędziami a węzłami.

Architektury takie jak MillenniumDB stosują dodatkowe algorytmy planowania zapytań tak, aby minimalizować konieczność „przekładania” powiązań w formacie trójek RDF na model grafu domenowego podczas wykonywania zapytania. Jedną ze strategii jest kompilacja wzorca do automatu analizującego relacje w locie – pozwala to uniknąć konieczności materializacji wszystkich potencjalnych ścieżek przed selekcją wyników. W kontekście optymalizacji przechodzenia istotna jest także kontrola długości ścieżek przeszukiwania. Składnia Cypher typu '-[:REL*1..3]->' umożliwia określenie minimalnej i maksymalnej liczby krawędzi pomiędzy punktami początkowym i końcowym, co ogranicza eksplorację grafu wyłącznie do zakresu istotnego dla analizy. W Gremlinie podobną funkcjonalność zapewnia '.repeat(...).times(n)' oraz '.until(...)', które dają większe możliwości dynamicznego kończenia traversała zgodnie z warunkami domenowymi. W obu przypadkach chodzi o uniknięcie zbędnego przemierzania fragmentów grafu, które nie mają wpływu na końcowy wynik.

Zoptymalizowane zapytanie powinno również wykorzystywać agregacje w taki sposób, aby działały lokalnie na ograniczonym zbiorze danych zamiast globalnie na całym grafie. Mechanizmy grupowania ('GROUP BY' w Cypher analogiczne do '.group()' w Gremlin) można stosować po wcześniejszym zawężeniu zbioru wejściowego – zmniejsza to zarówno wykorzystanie pamięci operacyjnej, jak i czas konieczny do przetworzenia wyników. Nadmierne korzystanie z agregacji przed filtracją prowadzi do niepotrzebnego powiększenia przestrzeni roboczej silnika. Konfiguracje systemowe odgrywają tu dużą rolę: baza wspierająca ACID często implementuje blokady transakcyjne zapewniające spójność zapisów kosztem wydajności odczytu. Strategia kłastrowania master-slave wymaga odpowiedniego rozdzielenia roli serwerów obsługujących zapis od tych, które realizują intensywne operacje transakcyjne lub analityczne. W środowiskach OLTP optymalizacja planów wykonania jej powinna być skorelowana z fizycznym rozmieszczeniem fragmentów grafu (sharding) tak, aby minimalizować liczbę przekroczeń granic partycji podczas traversała. Optymalizacja dotyczy także wyboru języka zapytań adekwatnego do scenariusza analitycznego – SPARQL wymaga rekonstrukcji wielu trójek przy każdej kwerendzie, co czyni go mniej efektywnym przy głębokich analizach strukturalnych niż Cypher, który interpretuje relację jako pojedynczy rekord zawierający pełen kontekst. Z kolei Gremlin sprawdza się tam, gdzie konieczne jest proceduralne



sterowanie ścieżką oraz dynamiczne podejmowanie decyzji o dalszych krokach przechodzenia; jednak brak deklaratywnej optymalizacji przez planowanie powoduje, że odpowiedzialność za wydajność spoczywa w większym stopniu na autorze zapytania.

Na poziomie modelowania schematu można wdrażać zasady projektowe usprawniające późniejsze wykonywanie zapytań, m.in. dobór kierunkowości krawędzi zgodnie z typowymi wzorcami wyszukiwania czy eliminację nadmiarowych właściwości utrudniających indeksowanie. Odpowiednie przypisanie etykiet do elementów pozwala stosować selektywne indeksy lokalne oraz szybkie filtry semantyczne w pierwszych krokach kwerendy. Wreszcie należy brać pod uwagę fizyczny mechanizm przechowywania – np. Neo4j utrzymuje własne pliki dla właściwości (`propertystore.db`) gdzie ich układ wpływa na szybkość odczytu sekwencyjnego. Świadomość tego układu oraz unikanie fragmentacji danych o wysokiej częstotliwości dostępu to elementy optymalizacji wymagające współpracy zespołu projektowego z administracją bazy. Całość procesu można sprowadzić do kilku powtarzalnych zasad: filtrować wcześniej i selektywnie; wybierać punkty startowe przechodzenie według dostępnych indeksów; ograniczać długość ścieżek; grupować dane dopiero po redukcji zbioru; unikać nadmiarowych joinów między strukturami; dostosować kierunki relacji i układ fizyczny danych do typowych scenariuszy eksploracyjnych; oraz dobierać język zapytań zgodny z oczekiwanym sposobem analizy grafu. Świadome ich stosowanie pozwala przekształcić nawet dużą i gęsto połączoną strukturę grafową w środowisko oferujące szybkie odpowiedzi bez kompromisu jakości wyników.

6.2 Przechowywanie i indeksowanie danych

Efektywność baz grafowych w dużej mierze zależy od sposobu, w jaki przechowywane są dane oraz jak zorganizowane jest ich indeksowanie. W systemach opartych na natywnym przechowywaniu grafów każdy węzeł i każda krawędź utrzymują bezpośrednie referencje do elementów powiązanych. Mechanizm taki – znany jako sąsiedztwo bez indeksowe – pozwala na przechodzenie po relacjach bez dodatkowych operacji wyszukiwania globalnego, co skraca czas odpowiedzi przy zapytaniach wymagających wielu skoków po grafie. Strukturalna organizacja danych przewiduje osobne pliki lub segmenty pamięci dla wierzchołków, krawędzi oraz właściwości; w Neo4j np. magazyn `'propertystore.db'` przechowuje atrybuty w formie list rekordów połączonych wskaźnikami. Dzięki takiemu podziałowi odczyt atrybutów czy relacji może być realizowany sekwencyjnie i z pominięciem fragmentów modelu nieistotnych dla bieżącego zapytania.



Magazyny nienatywne, które traktują graf jako nadbudowę nad innym formatem (np. tabelarycznym), muszą rekonstruować powiązania za każdym razem na podstawie kluczy obcych lub identyfikatorów przechowywanych wewnątrz rekordów. W praktyce oznacza to wzrost kosztu przeszukiwania, globalny indeks staje się konieczny do zlokalizowania sąsiadów danego wierzchołka, a każde dodatkowe przejście wymaga osobnego zapytania do warstwy magazynującej. Zastosowanie systemu natywnego likwiduje tę barierę; mechanizm indeksowania ogranicza się wtedy zwykle do wyszukiwania punktu startowego kwerendy. Rola indeksów pozostaje jednak istotna przy operacjach punktowych, gdy trzeba odnaleźć jeden lub kilka konkretnych elementów na podstawie wartości właściwości. Cypher wspiera tworzenie indeksów dla kombinacji etykiety i właściwości ('CREATE INDEX ON:Venue(name)'), a także unikalnych ograniczeń ('CREATE CONSTRAINT ON (c:Country) ASSERT c.name IS UNIQUE'). Indeksy te pozwalają przyspieszyć start przechodzenia – uwagę skupia się od razu na zbiorze początkowym zawężonym według kryterium biznesowego.

W scenariuszach intensywnego dodawania i modyfikowania danych należy rozważyć kompromis między liczbą utrzymywanych indeksów a kosztami ich aktualizacji. Wydajne środowisko grafowe może stosować także specjalizowane formy indeksowania oparte na strukturze danych grafu. Rozwiązania takie jak gIndex wykorzystują częste podstruktury jako element encji indeksowej. Dzięki temu powtarzalne wzorce ścieżek mogą być znajdowane szybciej niż poprzez standardowe przechodzenie. Tego rodzaju metoda jest szczególnie efektywna przy analizach offline dużych baz trójek RDF, gdzie wyszukanie konkretnego wzorca trójki wymaga filtrowania spośród miliardów rekordów.

Przechowywanie metadanych stanowi kolejny obszar wpływający na wydajność. W grafie właściwości zarówno węzły, jak i krawędzie mają przypisane zestawy właściwości; ich fizyczna lokalizacja względem komponentu bazowego ma znaczenie przy dostępie sekwencyjnym. System może umieszczać często używane atrybuty w blokach o stałej długości lub korzystać z kodowania binarnego skracającego czas deserializacji. Nadmierna liczba właściwości zwiększa rozmiar rekordu i może wydłużyć odczyt, stąd projektowo warto oddzielić metadane o niskiej częstotliwości użycia. Aspekt wysokiej dostępności wymaga uwzględnienia technik replikacji i rozproszenia magazynu grafowego, co pociąga za sobą kwestię dystrybucji wskaźników przylegania. Fragmentowanie grafu wiąże się z ryzykiem konieczności przekraczania granic partycji podczas przechodzenia – optymalny podział powinien minimalizować liczbę takich operacji. Przy tym zastosowanie replikowanych fragmentów dla najczęściej odwiedzanych węzłów ułatwia równoważenie obciążenia. W systemach hybrydowych integrujących graf właściwości z RDF konieczne jest utrzymywanie spójności pomiędzy dwoma formami zapisu relacji. Grafy domenowe



zachowują kompatybilność poprzez zastosowanie wspólnego modelu rdzeniowego oraz identyfikatorów krawędzi (ID krawędzi). Takie rozwiązanie wymaga odrębnej warstwy mapowania podczas zapisu i odczytu, indeksy muszą obsługiwać nie tylko tradycyjne filtry po właściwościach, ale też szybkie tłumaczenie predykatów RDF na typy relacji grafu właściwości. Niektóre bazy stosują również konstrukcje typu krawędzie wirtualne, które istnieją jedynie logicznie i są generowane podczas wykonywania zapytań na podstawie reguł dotyczących właściwości elementów. Choć zmniejsza to zużycie pamięci stałej, nakłada dodatkowe wymagania na proces planowania kwerendy, generowanie takich połączeń musi być powiązane z mechanizmami filtrowania wyników, by uniknąć eksplorowania zbędnych części struktury.

Od strony zarządzania integralnością danych przechowywanie informacji w natywnym formacie grafowym pozwala wdrożyć pełen model ACID dla transakcji. Zapewnia to spójność nawet przy równoczesnym dostępie wielu instancji klastra. Przy architekturach BASE nastawionych na skalowalność horyzontalną magazyn może dopuszczać opóźnioną spójność danych – wpływa to jednak na gwarancje dotyczące poprawności wyników zapytań bezpośrednio po modyfikacjach. Podsumowując kwestie przechowywania i indeksowania można wskazać kilka kluczowych zaleceń projektowych wynikających z doświadczeń praktycznych: wykorzystanie sąsiedztwa bez indeksowego do minimalizacji kosztu przechodzenia; tworzenie precyzyjnych indeksów tylko dla często używanych punktów startowych; wyważenie liczby utrzymywanych metadanych względem kosztu ich odczytu; zastosowanie strukturalnych metod indeksowania dla powtarzalnych wzorców; unikanie przeciążenia przez nadmiar mało użytecznych właściwości; oraz dopasowanie strategii przechowywania do wymogów dostępności i spójności specyficznych dla danego kontekstu aplikacyjnego. Tylko takie całościowe podejście gwarantuje optymalne wykorzystanie potencjału technologii baz grafowych w zadaniach analitycznych i transakcyjnych.

7 Zastosowania baz grafowych

7.1 Sieci społecznościowe

Analiza połączeń między użytkownikami

Analiza połączeń między użytkownikami w bazach grafowych opiera się na przeglądaniu i interpretacji sieci relacji jako dynamicznego układu węzłów reprezentujących osoby oraz krawędzi odwzorowujących różne typy interakcji. Kluczowym elementem jest tu możliwość bezpośredniego przechodzenia od jednego



węzła do powiązanych partnerów poprzez mechanizm sąsiedztwa bez indeksowego, co eliminuje konieczność wykonywania kosztownych operacji wyszukiwania w globalnych strukturach. Takie podejście pozwala uzyskać obraz wzajemnych powiązań niemal natychmiast po wskazaniu punktu startowego – istotne, gdy analizowane są gęsto splecione sieci społecznościowe z dużą liczbą relacji per użytkownik. Podczas modelowania analizy połączeń warto uwzględnić charakter relacji – mogą być one jednokierunkowe, jak obserwowanie na platformach mikroblogowych lub symetryczne, jak znajomość w serwisach społecznościowych. Kierunkowość determinuje zarówno logikę zapytań, jak i wykorzystywane algorytmy – wyszukiwanie osób obserwowanych przez danego użytkownika odbywa się przy przechodzeniu wychodzących krawędzi, podczas gdy identyfikacja wspólnych znajomych wymaga sprawdzenia przecięcia zestawów sąsiadów z obu stron. W praktyce często powstaje potrzeba analizy "friend-of-a-friend", która polega na przechodzeniu ścieżki o długości dwóch krawędzi danego typu; przykładowo '(u1)-[:FRIEND]->(u2)-[:FRIEND]->(u3)' służy identyfikacji potencjalnych powiązań niejawnych.

Bazy grafowe wspierają tego typu scenariusze dzięki ekspresji języków zapytań takich jak Cypher, które umożliwiają proste definiowanie wzorców strukturalnych i ich filtrowanie według właściwości użytkowników czy samych relacji. Można np. wskazać wyłącznie połączenia utworzone w ostatnich 30 dniach lub relacje o wysokiej intensywności interakcji – parametr ten zapisany jest jako właściwość krawędzi 'interactionFrequency'. W serwisie Talent.net przykład rekomendacji zawodowych opierał się na wzorcu wspólnych zainteresowań: znalezienie połączeń pomiędzy Sarah i Benem z dwoma wspólnymi interesariuszami dawało silniejsze skojarzenie niż jedno podobieństwo z Charliem. Siła więzi może być tu modelowana liczbowo jako cecha krawędzi i używana do sortowania wyników wyszukiwanych kontaktów. Analiza połączeń obejmuje także wykrywanie struktur lokalnie spójnych, takich jak kliki czy społeczności tematyczne. Algorytmy klasteryzujące, uruchamiane na pełnym zbiorze danych grafowych, pozwalają segmentować sieć według gęstości powiązań. Wykorzystanie semantycznie wzbogaconych metadanych dodatkowo podnosi precyzję takich działań – przykładowo etykietowanie użytkowników według branży lub lokalizacji tworzy warstwę kontekstową dla wyników.

Wiedza o lokalizacji może być sprzężona z innymi źródłami danych – systemy głębszej analizy tekstu potrafią rozstrzygać niejednoznaczne nazwy miejsc czy osób i wiązać je z odpowiednimi profilami w grafie, co usprawnia identyfikację realnych połączeń przypisanych do tego samego bytu. Interpretacja sieci połączeń obejmuje sytuacje z aliasowaniem tożsamości. W detekcji nadużyć może wystąpić przypadek kilku kont prowadzonych przez jedną osobę lub używanie pseudonimów do



kreowania pozornych rozmów. Analiza tych wzorców wymaga śledzenia grup krawędzi między zestawem powiązanych węzłów oraz porównania ich struktury komunikacyjnej z typowymi zachowaniami użytkowników niebiorących udziału w procederze. Dzięki możliwości traktowania relacji jako odrębnych rekordów z własnymi właściwościami możliwe jest oznaczenie podejrzanej interakcji flagą lub wagą ryzyka.

Część analiz koncentruje się na identyfikowaniu punktów centralnych sieci – tak zwanych hubów. Wskaźniki takie jak stopień wierzchołka, wyśrodkowanie pomiędzy czy bliskość umożliwiają wskazanie użytkowników posiadających największą liczbę unikalnych połączeń lub kontrolujących największą liczbę ścieżek przepływu informacji. Zbudowanie rankingu wyśrodkowania pozwala lepiej zrozumieć strukturę wpływów i może służyć do optymalizacji kampanii marketingowych czy wykrywania grup liderów opinii. Obok klasycznej eksploracji topologicznej wartościowe są zapytania semantyczne integrujące różne rodzaje relacji. Można łączyć dane o kontaktach zawodowych ze wspólnym uczestnictwem w wydarzeniach – model krawędzi jako węzeł pozwala wtedy potraktować wydarzenie jako osobny wierzchołek powiązany krawędziami do wszystkich uczestników. Taki schemat ułatwia analizę nakładania się różnych aspektów aktywności społecznej.

Warto zauważyć rolę wizualizacji wyników analizy – choć bazy grafowe same nie wymagają prezentowania danych jako grafu, narzędzia wizualizacyjne pomagają szybko uchwycić rozbudowane zależności pomiędzy użytkownikami oraz zauważyć anomalie takie jak odizolowane grupy lub nienaturalnie gęste segmenty sieci. Wizualna interpretacja wsparta możliwością interakcji z modelem zwiększa efektywność pracy badacza. Łączenie omawianych metod analizy połączeń z innymi wspomnianymi wcześniej technologiami pozwala uzyskać środowisko zdolne do błyskawicznej obsługi skomplikowanych zapytań nawet dla ogromnych sieci użytkowników. Świadome wykorzystanie struktury grafowej, właściwości relacji oraz dodatkowych źródeł metadanych stanowi o skuteczności tych działań oraz jakości uzyskiwanych wyników.

Wykrywanie społeczności

Wykrywanie społeczności w grafach będących odwzorowaniem sieci społecznych polega na identyfikacji skupisk węzłów o wyższej gęstości połączeń wewnętrznych niż zewnętrznych. Proces ten odwołuje się do analizy struktury topologicznej oraz atrybutów relacji i węzłów, które mogą wskazywać na naturalne grupowanie obiektów. W kontekście baz grafowych kluczowe jest wykorzystanie mechanizmu sąsiedztwa bez indeksowego – każda krawędź posiada odniesienia do sąsiadujących węzłów, co pozwala błyskawicznie przechodzić pomiędzy elementami



i obliczać miary spójności czy centralności bez dodatkowego odwoływania się do indeksów globalnych. Dzięki temu algorytmy detekcji społeczności mogą działać w trybie interaktywnym nawet przy dużych strukturach.

Podstawą wielu metod wykrywania społeczności jest analiza macierzy sąsiedztwa lub list sąsiadów dla każdego wierzchołka, przy czym warstwa magazynowa grafu właściwości ułatwia pobieranie zarówno topologii połączeń, jak i związanych z nimi właściwości opisowych. Możliwe jest uwzględnianie dodatkowych metadanych, takich jak częstotliwość interakcji czy typ powiązania, co pozwala różnicować wagę krawędzi w modelu grafowym. W praktyce stosuje się heurystyki maksymalizacji modularności, gdzie celem jest uzyskanie podziału z wysokimi wartościami tej miary, oznaczającej przewagę liczby połączeń wewnątrz społeczności nad spodziewaną liczbą połączeń losowych. Mechanizm grafu właściwości pozwala te obliczenia przeprowadzać bezpośrednio na danych opatrzonych etykietami i właściwościami. Integracja wykrywania społeczności z językami zapytań takimi jak Cypher umożliwia implementację filtrów wyjściowych, po uruchomieniu algorytmu można wydzielać podgrupy według dodatkowych kryteriów biznesowych lub semantycznych. Na przykład `'MATCH (p:Person)-[:FRIEND]->(f) WHERE p.sector = 'IT'` pozwala ograniczyć analizowaną podstrukturę do osób związanych zawodowo z technologiami. Takie selekcje można łączyć z wynikami algorytmów Louvain czy Propagacji etykiet uruchamianych jako procedury systemowe.

Ważnym aspektem jest również wykorzystywanie krawędzi jako węzeł przy modelowaniu relacji uczestnictwa w wydarzeniach; umożliwia to identyfikację społeczności bazującej na współuczestnictwie zamiast bezpośrednich kontaktach. W takim przypadku każdy wierzchołek–wydarzenie łączy krawędzie do uczestników, a detekcja gęstych subgrafów wokół tych elementów przechwytuje społeczności tematyczne. Techniki oparte na analizie ścieżek wykorzystują rozszerzone mechanizmy przechodzenia do badania przepływu informacji między węzłami. Gremlin oferuje możliwość iteracyjnego budowania zbiorów odwiedzonych węzłów (`'repeat(...).times(n)'`) oraz dynamicznego zatrzymywania przechodzenia po spełnieniu określonego warunku `'until(...)'`. Dzięki temu można badać zmienne głębokości kontaktów i wychwytywać granice naturalnych skupisk powiązań. W połączeniu z krokami agregującymi `'group()'` czy `'groupCount()'` możliwe jest tworzenie map odwzorowujących rozkład liczebności powiązań wewnątrz kandydujących społeczności.

Przechowywanie wskaźników wyśrodkowania i wyników wcześniejszych analiz jako właściwości węzłów lub krawędzi upraszcza późniejsze filtrowanie rezultatów detekcji. Na przykład wartość `communityId` przypisana każdemu węzłowi może posłużyć do szybkiego wyodrębnienia całej grupy poprzez prosty filtr etykiety i



właściwości. Podobnie przechowywanie wyników hierarchicznego grupowania umożliwia ruchome definiowanie poziomu szczegółowości wykrytych struktur. W środowiskach wymagających integracji danych heterogenicznych, takich jak trójki RDF połączone z grafem właściwości, wykrywanie społeczności może korzystać ze wspólnego rdzenia semantycznego grafu domenowego. Predykaty RDF odwzorowane na typy relacji grafu właściwości pozwalają używać jednolitych algorytmów niezależnie od źródła danych. Algorytmy operujące na trójkach (np. SPARQL Group By) mogą być tłumaczone na konstrukcje przechodzenia Gremlina bądź wzorce Cyphera.

Detekcja społeczności znajduje zastosowanie m.in. przy rekomendacjach treści opartych o interesy grupowe, identyfikacji klastrow podejrzanej aktywności czy analizie wpływu liderów opinii. Przykład praktyczny: jeśli użytkownik A uczestniczył wraz z kilkoma członkami istniejącej społeczności w tym samym wydarzeniu oraz posiada bezpośrednie powiązania przez co najmniej dwóch z nich, system może zakwalifikować go jako potencjalnego nowego członka grupy. Takie reguły można implementować poprzez sekwencje zapytań agregujących liczby wspólnych relacji i filtrujących według progowej wartości tej metryki. Warto też wspomnieć o roli wizualizacji wyników wykrywania społeczności, graficzne reprezentacje klastrow pomagają zweryfikować poprawność działania algorytmów i dostrzec anomalie takie jak nienaturalnie duże lub zwarte grupy powiązań. Połączenie obrazu topologii z metadanymi opisowymi elementów ułatwia interpretację kontekstu badanego skupiska. Cały proces skutecznego wykrywania społeczności wymaga więc spójnego modelu danych umożliwiającego szybkie przechodzenie, bogatego zestawu właściwości elementów struktury oraz elastycznych języków zapytań integrujących wyniki obliczeń z logiką aplikacyjną systemu grafowego.

7.2 Systemy rekomendacyjne

W kontekście zastosowań grafowych baz danych systemy rekomendacyjne korzystają z możliwości modelowania bezpośrednich i pośrednich relacji w celu przewidywania preferencji użytkowników oraz sugerowania im zasobów potencjalnie interesujących. Fundamentem ich działania jest analiza struktury powiązań, gdzie węzły mogą reprezentować osoby, produkty, treści medialne czy usługi, a krawędzie opisują interakcje takie jak zakup, polubienie, recenzja lub wspólne uczestnictwo w wydarzeniu. W odróżnieniu od klasycznych mechanizmów wyszukiwawczych operujących głównie na dopasowaniu słów kluczowych, tutaj rdzeniem logiki jest wykrywanie wzorców relacyjnych, identyfikacja tego, z czym użytkownik miał kontakt, jak często i w jaki sposób. Wyraźne relacje powstają tam, gdzie użytkownik sam



deklaruje powiązanie, np. oznacza produkt jako ulubiony lub dodaje innego użytkownika do grona znajomych. Z kolei relacje ukryte są wnioskowane na podstawie działań powiązanych pośrednio: jeżeli dwóch różnych klientów kupuje ten sam zestaw produktów lub odwiedza te same lokalizacje online, system może uznać ich preferencje za częściowo zgodne. Ta druga kategoria stanowi bogate źródło informacji przy predykcjach, ujawnia skryte zależności niewidoczne wprost.

Algorytmy rekomendacyjne mają charakter indukcyjny i sugestywny, co oznacza, że nie dążą do dowodu ścisłej poprawności wyniku, lecz do wyłonienia zbioru przypuszczalnie istotnych propozycji na bazie podobieństwa zachowań. Obejmuje to zarówno wyszukiwanie elementów atrakcyjnych dla określonego użytkownika, jak i wskazywanie grup odbiorców zainteresowanych konkretnym zasobem. Model grafowy upraszcza ten proces dzięki temu, że konkretna interakcja jest pełnoprawnym zasobem posiadającym własne właściwości, przykładowo krawędź RATES pomiędzy klientem a produktem może zawierać ocenę liczbową oraz datę wystawienia recenzji. W praktyce system rekomendujący może działać iteracyjnie w oparciu o przechodzenie: startując od węzła reprezentującego użytkownika, przechodzi do produktów ocenionych wysoko przez niego lub jego bezpośrednich znajomych i szuka dalszych zasobów powiązanych z tymi produktami.

Mechanizm sąsiedztwa bez indeksowego pozwala wykonywać te kolejne kroki niemal natychmiast dzięki fizycznym wskaźnikom między rekordami. Tak budowana ścieżka uwzględnia również wagę relacji, krawędzie mogą być filtrowane według progu oceny lub częstotliwości interakcji, co zmniejsza ryzyko generowania przypadkowych rekomendacji. Rozszerzeniem podstawowej logiki jest integracja wiedzy strukturalnej z dodatkowymi danymi semantycznymi przechowywanymi w grafach wiedzy. Połączenie obiektów biznesowych (produktów, usług) z ontologiami pozwala opisać atrybuty kategorii czy funkcji danego zasobu niezależnie od konkretnego egzemplarza. Dzięki temu można tworzyć rekomendacje wielopoziomowe, sugerować produkt nie tylko dlatego, że inni go wybrali, ale także z powodu podobieństwa jego cech technicznych do czegoś już posiadanego przez użytkownika.

Systemy grafowe wspierają też scenariusze inferencji: mechanizm reguł logicznych może automatycznie dodawać nowe relacje wynikające ze znanych faktów. Na przykład, jeśli krawędzie wskazują na częste zakupy produktu X wraz z Y oraz produktu Y wraz z Z, silnik może zasugerować X+Z mimo braku ich bezpośredniej relacji zakupowej przez danego klienta. W praktyce przekłada się to na zwiększenie zakresu "znanych" preferencji bez konieczności gromadzenia pełnej historii zachowań. W projektowaniu schematu grafu pod kątem rekomendacji istotna jest decyzja o modelowaniu wielowymiarowych powiązań jako krawędź jako węzeł.



Węzeł reprezentujący zdarzenie zakupu może mieć krawędzie do klienta, produktu oraz kampanii marketingowej; taki układ pozwala łatwo śledzić efektywność kampanii względem konkretnych grup docelowych i współdzielić dane między różnymi modułami analitycznymi. Analiza takich meta-relacji uzupełnia tradycyjne filtry treściowe.

Wzorce zapytań dla systemów rekomendacyjnych często wykorzystują zmienną długość ścieżki. Chodzi o uchwycenie bardziej złożonych asocjacji niż prosta więź typu klient–produkt; przykładowo '(user)-[:PURCHASED]->(item)<-[:PURCHASED]-(otherUser)-[:VIEWED]->(newItem)'. Tutaj 'newItem' staje się kandydatem do polecenia pierwszemu użytkownikowi, mimo że brak bezpośredniej historii interakcji. Wykrycie grupy użytkowników o silnych więziach wewnętrznych daje dodatkowy kontekst: jeśli większość członków społeczności zaczyna używać nowej usługi, prawdopodobieństwo jej akceptacji przez pozostałych rośnie. Można zatem traktować wykryte klastry jako jednostki rekomendacyjne i kierować spersonalizowane oferty nie do pojedynczych osób, lecz całych mikrospołeczności.

Technologicznie wdrożenie takich systemów wymaga równowagi między wydajnością a elastycznością modelu danych. Rozwiązania natywne zapewniają minimalny czas reakcji przy przechodzeniach wielostopniowych; hybrydowe podejścia grafu własności i domeny dają interoperacyjność ze źródłami RDF używającymi standardów sieci semantycznej. Dobór języka zapytań także ma znaczenie: Cypher sprawdza się przy prostych wzorcach typu "klient → produkt → inny klient → inny produkt", Gremlin może być konieczny przy bardziej proceduralnym sterowaniu długością i warunkami przeszukiwania podgrafu. Nie można pominąć roli optymalizacji indeksowania – selektywne indeksy dla właściwości takich jak category czy tag redukują czas startu zapytań filtrujących kandydatów rekomendacji po cechach domenowych. W jeszcze bardziej dynamicznych środowiskach stosuje się też cache wyników popularnych przechodzeń oraz przed obliczeniowe rankingi podobieństw użytkownik–produkt dla szybkiej prezentacji sugestii. Podsumowując techniczny aspekt implementacji – skuteczny system rekomendacyjny oparty na bazie grafowej powinien potrafić łączyć różnorodne typy relacji jawnej i ukrytej, korzystać z bogactwa metadanych opisujących zarówno obiekty jak i same więzi między nimi, integrować inferencję oraz mechanizmy krawędzi jako węzeł dla analizy meta-interakcji, a także adaptować strategię przechodzenia do specyfiki danego przypadku biznesowego czy badawczego. Tylko spójne wykorzystanie tych elementów gwarantuje wysoką trafność sugerowanych zasobów i efektywność ich dostarczania w czasie rzeczywistym.



7.3 Analiza sieci powiązań biznesowych

Analiza sieci powiązań biznesowych przy użyciu baz grafowych opiera się na modelowaniu relacji między podmiotami gospodarczymi – zarówno tych bezpośrednich, jak kontrakty, partnerstwa czy transakcje, jak i pośrednich wynikających ze struktury rynku lub historii współpracy. Węzły mogą reprezentować firmy, jednostki organizacyjne, pracowników, dostawców lub klientów, a krawędzie odwzorowują różnorodne więzi: od dostaw łańcucha logistycznego po współdzielenie projektów badawczo-rozwojowych (R&D). Podejście grafowe umożliwia analizę złożonych układów zależności w sposób trudny do realizacji na schematach tabelarycznych, ponieważ relacje są traktowane jako obiekty pierwszego rzędu z własnymi właściwościami. Dzięki temu można zapisywać na poziomie krawędzi np. wartość finansową kontraktu, datę rozpoczęcia współpracy czy wskaźniki jakości partnerstwa. Przydatność takiego modelu wzrasta w kontekście tzw. Business 360 – koncepcji integracji wszystkich istotnych danych o interakcjach firmy z interesariuszami w jeden spójny obraz.

Tworzenie pełnej perspektywy wymaga zebrania rozproszonych informacji z działów sprzedaży, marketingu, finansów i obsługi klienta oraz ich połączenia według powiązań między jednostkami organizacyjnymi i podmiotami zewnętrznymi. Bazy grafowe pozwalają na skorelowanie faktów wskazujących np. ostatnią interakcję z konkretnym dostawcą, co niweluje potrzebę przeszukiwania oddzielnych systemów. Mechanizm sąsiedztwa bez indeksowego sprawia, że ustalenie ciągu zdarzeń między dwoma podmiotami odbywa się w czasie rzeczywistym nawet przy setkach tysięcy relacji. Kluczowym elementem analizy sieci biznesowej jest mapowanie struktur własności i zależności kapitałowych. Węzeł odpowiadający przedsiębiorstwu może być powiązany wieloma krawędziami typu OWNS, HAS_SHARE_IN czy IS_SUBSIDIARY_OF do innych firm. Dzięki temu możliwe staje się ujawnianie powiązań korporacyjnych blokujących konkurencję lub tworzących oligopol na danym rynku. Analiza ścieżek o zmiennej długości pozwala dotrzeć od klienta końcowego do właściciela dominującego w kilku krokach nawet w strukturach rozciągniętych geograficznie.

Bazy grafowe ułatwiają także odwzorowanie łańcuchów dostaw – zwykle obejmujących wielu pośredników oraz operacje transgraniczne. Wykorzystując etykiety krawędzi i właściwości opisujące parametry logistyczne (czas realizacji, koszt transportu), można analizować przepływy towarów i identyfikować wąskie gardła procesu. Geograficzny wymiar analizy zostaje zachowany dzięki możliwości przypisania lokalizacji zarówno dostawcom jak i transakcjom, relacje mogą mieć atrybut określający punkt nadania czy przeładunku. Integracja danych z różnych



domen poprzez modele między domenowe pozwala uchwycić efekty sieciowe w całym łańcuchu wartości. Modelowanie „grafu grafów” umożliwia połączenie subdomen – np. produkcji, dystrybucji i sprzedaży – bez utraty szczegółowości właściwej każdej z nich. Ta technika jest szczególnie cenna przy analizach strategicznych, przedsiębiorstwo może ocenić wpływ zmiany jednego dostawcy komponentu na czas realizacji zamówień końcowych u kluczowych klientów. Wykorzystanie grafów wiedzy semantycznej poszerza zakres możliwych analiz o powiązania semantyczne dotyczące obiektów biznesowych. Różne identyfikatory, nazwy handlowe oraz role danego podmiotu są konsolidowane w jednym wierzchołku opisanym bogatym zestawem właściwości. Pozwala to stworzyć widok 360 stopni dla klientów lub partnerów zawierający nie tylko transakcje i płatności, ale też informacje o preferencjach zakupowych bądź historii kontaktów. Tak integracja danych ustrukturyzowanych i nieustrukturyzowanych umożliwia systemowi rekomendowanie strategicznych działań, np. ofert sprzedaży krzyżowych opartych na analizie poprzednich współprac.

Od strony implementacyjnej języki zapytań jak Cypher umożliwiają formułowanie wzorców opisujących interakcje między podmiotami takie jak '(c1:Company)-[r:PARTNER_OF]->(c2:Company)' filtrowane względem wartości kontraktów zapisanych w 'r.value'. Filtrowanie wyników według przedziału czasowego czy sektora działalności zawartego we właściwościach węzłów daje narzędzie do badania trendów rynkowych, np. rosnącej intensywności współpracy pomiędzy przedsiębiorstwami technologicznymi a instytucjami finansowymi. Analizy wykrywania społeczności przedstawione wcześniej znajdują naturalne zastosowanie również tutaj – grupy przedsiębiorstw o wysokiej gęstości relacji mogą odpowiadać realnym klastrom gospodarczym bądź stowarzyszeniom branżowym. Identyfikacja takich klastrow pozwala ocenić siłę kooperacyjną danego regionu lub branży oraz przewidzieć kierunki rozwoju przez obserwację przepływu zasobów pomiędzy członkami grupy.

Cennym polem zastosowania jest analiza ryzyka w sieciach powiązań biznesowych – mapując relacje kredytowe czy gwarancyjne można wykrywać potencjalne punkty awarii sieci finansowej przedsiębiorstw. W przypadku utraty płynności przez jedno ogniwo kaskadowe skutki mogą objąć dużą część sieci zależnej; symulacja takich scenariuszy odbywa się poprzez iteracyjne usuwanie krawędzi lub zmianę ich wag odpowiadających wypłacalności podmiotów. Łączenie warstw informacyjnych, danych teleadresowych, dokumentacji prawnej, historii zatrudnienia kadry zarządzającej, buduje wielowymiarowy graf relacji biznesowych o dużej wartości analitycznej dla audytorów czy inwestorów. Możliwość operowania równocześnie na



topologii połączeń oraz szczegółowych metadanych zwiększa precyzję prognoz dotyczących zachowań rynkowych badanego ekosystemu.

Ostatecznie analiza sieci powiązań biznesowych wspierana przez technologie grafowe daje narzędzia do mapowania pełnego obrazu interakcji gospodarczych; umożliwia zarówno operacyjne zarządzanie bieżącymi procesami (monitoring kontraktów, optymalizacja dostaw) jak też strategiczne planowanie rozwoju przez modelowanie wpływu nowych relacji lub likwidacji istniejących na całą sieć partnerstw. Kombinacja wydajnej architektury natywnej z elastycznym modelem danych pozwala reagować na zmiany rynkowe szybciej niż przy wykorzystaniu hermetycznych systemów relacyjnych, które wymagają żmudnego scalania informacji przed każdą głębszą analizą.

7.4 Bezpieczeństwo i wykrywanie oszustw

Analiza wzorców oszustw

Wykrywanie i analiza wzorców oszustw w grafowych bazach danych opiera się na możliwości modelowania relacji w sposób pozwalający śledzić zarówno bezpośrednie interakcje, jak i ukryte powiązania pomiędzy obiektami. Każda wiadomość, transakcja czy zdarzenie może być formalnie odwzorowane jako odrębny węzeł połączony krawędziami z uczestnikami, umożliwia to rejestrowanie kontekstu zdarzenia oraz dodatkowych właściwości technicznych lub semantycznych. Takie podejście pozwala uniknąć uproszczonego modelu, gdzie operacja jest jedynie etykietą na relacji; dzięki centralnemu umieszczeniu „wiadomości” w strukturze grafowej można analizować jej powiązania z wieloma odbiorcami, nadawcami bądź kanałami komunikacyjnymi. Mechanizm sąsiedztwa bez indeksowego, typowy dla systemów natywnych, daje możliwość przeszukiwania relacji z wysoką efektywnością, w przypadku analizy wykrywania oszustw jest to kluczowe przy scenariuszach obejmujących przeskanowanie tysięcy powiązań w celu znalezienia nietypowych wzorców.

Przykładowym zastosowaniem jest wykrywanie sytuacji, gdy użytkownik wysłał wiadomość na swój własny adres aliasowy przez pole CC. Powiązanie dwóch krawędzi, SENT oraz CC, ze wspólnym kontekstem aliasu można łatwo wyrazić jako wzorzec zapytania Cyphera o strukturze kilku hopów; taki mechanizm pozwala szybko wyłowić anomalie z dużego zbioru danych. Podobna logika znajduje zastosowanie w analizie przepływów finansowych, poprzez reprezentację każdej transakcji jako oddzielnego węzła można badać kierunki jej przepływu, atrybuty płatności i powiązania rachunków bankowych. Zastosowanie filtrów według czasu



trwania lub kwoty umożliwia wyłanianie grup operacji odbiegających od normy. W tym ujęciu wagę krawędzi można potraktować jako istotny parametr analizy: wysoka wartość finansowa lub krótki odstęp czasu między kolejnymi transferami często sygnalizuje działania próbujące obejść mechanizmy kontroli. Procedury detekcji realizowane są poprzez łączenie wzorców topologicznych z analizą właściwości elementów grafu. Można przykładowo zestawiać listę kont posiadających wspólnych dostawców usług pośrednich albo porównywać profile aktywności względem typowego zachowania danego segmentu użytkowników. Systemy grafowe wspierają tu agregacje oparte na klastrach relacji: wyszukiwanie grup kont współdzielących te same identyfikatory urządzeń lub adresy IP, co może oznaczać zorganizowaną siatkę oszustw.

Znaczenie ma także możliwość ewolucji schematu danych, dodanie nowego typu krawędzi opisującej specyficzny sygnał ostrzegawczy (np. `FAILED_LOGIN_AFTER_TRANSFER`) nie wymaga przebudowy całości bazy. Pozwala to reagować na pojawiające się techniki oszustw poprzez rozszerzenie modelu o nowe elementy monitoringu bez zakłócania istniejących procesów analitycznych. Struktury o wysokiej gęstości połączeń pomiędzy niewielką liczbą kont mogą wskazywać na kooperację w celu transferowania środków lub rozpraszania ich poza systemem kontroli. Algorytmy propagacji etykiet zastosowane do relacji transakcyjnych potrafią wydzielić takie skupiska automatycznie i przypisać im unikalne identyfikatory społeczności ryzyka.

Z technicznego punktu widzenia istotne jest dopasowanie języka zapytań do rodzaju wzorca poszukiwanego – Cypher nadaje się do prostych dopasowań typu "konto -> transakcja -> konto", Gremlin bywa wybierany przy konieczności proceduralnego sterowania długością przechodzenia oraz stosowania warunków zakończenia zależnych od wartości właściwości. W warunkach trójek RDF wykorzystuje się SPARQL do budowy zapytań operujących na predykatkach opisujących zdarzenia; hybrydowe grafy domenowe pozwalają tłumaczyć te predykaty na typy relacji grafu właściwości dla spójnej analizy.

Ważnym krokiem jest wdrożenie reguł inferencyjnych – zapisane fakty mogą generować nową wiedzę o potencjalnych zagrożeniach. Przykładem jest identyfikacja powtarzalnego cyklu transferów między tymi samymi podmiotami zakończonego zawsze wypłatą gotówki, reguła logiczna może dodać relację `SUSPECTED_CYCLE` do par uczestników bez manualnego przeglądania historii. Praktyczne wdrożenia pokazują także wartość integracji danych miękkich, jak historia logowań czy metadane urządzeń użytych podczas operacji (Author, b). Umożliwia to budowę bardziej kompletnych profili zachowań, które można porównywać do wzorców prawidłowej aktywności w celu wyznaczania poziomu ryzyka. Połączenia



tych informacji z rdzeniem relacyjnym grafu dają silne narzędzie detekcyjne zdolne reagować w czasie rzeczywistym na nietypowe działania. Efektywność takiej analizy zależy od jakości indeksowania punktów startowych zapytań i filtrowania wyników już we wczesnych etapach przechodzenia. System powinien kierować uwagę jedynie na elementy spełniające minimalne kryteria ryzyka, co ogranicza koszt obliczeniowy przy dużych wolumenach danych i skraca czas reakcji mechanizmów zabezpieczeń. Ostatecznie model grafowy zwiększa przejrzystość i skuteczność działań przeciwdziałania oszustwom dzięki integracji powiązań topologicznych, cech opisowych oraz dynamicznych reguł wykrywających coraz bardziej złożone formy nadużyć.

Identyfikacja podejrzanych połączeń

Identyfikacja podejrzanych połączeń w grafowych bazach danych jest procesem, który bazuje na szczegółowym odwzorowaniu relacji biznesowych, transakcyjnych lub komunikacyjnych jako węzłów i krawędzi posiadających własne właściwości oraz kontekst semantyczny. W modelu grafu właściwości każda relacja jest obiektem pierwszej klasy – wyposażonym w identyfikator, typ (etykietę) oraz zestaw atrybutów opisujących m.in. czas utworzenia, źródło danych czy parametry ilościowe takie jak wartość finansowa bądź liczba interakcji. Dzięki temu można prowadzić analizę nie tylko pod kątem samej topologii sieci połączeń, ale też ich jakości i okoliczności powstania. W praktycznej implementacji systemów bezpieczeństwa analiza takich połączeń wymaga określenia punktów startowych przechodzenia i wczesnego stosowania filtrów redukujących zbiór danych do elementów potencjalnie ryzykownych.

Mechanizm sąsiedztwa bez indeksowego pozwala na błyskawiczne przechodzenie od jednego wierzchołka do jego sąsiadów bez angażowania globalnych indeksów, co skraca czas reakcji przy dużej liczbie kaskadowych relacji. Filtry mogą opierać się o wartości właściwości krawędzi – np. oznaczenia typu highRiskFlag, nietypowe kanały komunikacji lub anomalne parametry transakcji. Przykładowo relacje TRANSFERRED_FUNDS o wysokiej kwocie wykonywane z rachunków o świeżym stażu mogą stanowić pierwsze sygnały ostrzegawcze. Istotnym wyzwaniem jest uchwycenie powiązań pośrednich między bytami, które na pierwszy rzut oka nie wykazują bezpośredniej interakcji. W tym celu stosuje się ścieżki o zmiennej długości – konstrukcje w Cypher `MATCH p=(a)-[*2...5]->(b)` lub analogiczne mechanizmy `.repeat(...).times(n)` w Gremlin umożliwiają badanie relacji wielostopniowych z uwzględnieniem ograniczeń głębokości.

W scenariuszach obejmujących np. pranie pieniędzy taka ścieżka może reprezentować sekwencję transferów przez konta pośredniczące należące do kilku



różnych podmiotów; dostrzeżenie wzorca „łańcuchowego” pozwala zaklasyfikować całą grupę powiązań jako podejrzaną. Przy bardziej zaawansowanych analizach warto wdrożyć koncepcję krawędzi jako węzeł dla meta-relacji opisujących samą więź. Jeśli relacja CONNECTION ma atrybuty wskazujące typ szyfrowania czy metodę autoryzacji użytkownika, można ją zamodelować jako osobny węzeł połączony z uczestnikami interakcji i innymi bytami kontekstowymi (np. urządzeniami końcowymi). Umożliwia to wyszukiwanie złożonych zależności między parametrami technicznymi połączenia a historią zachowań uczestników.

Systemy identyfikacji zagrożeń stosują reguły inferencyjne do tworzenia nowych relacji z już dostępnych danych. Możliwe jest np. automatyczne dodanie krawędzi SUSPECTED_LINK pomiędzy podmiotami, które biorą udział w przynajmniej trzech transakcjach powiązanych wspólnym adresem IP lub lokalizacją geograficzną zapisaną we właściwościach krawędzi. Takie reguły mogą działać cyklicznie lub w trybie czasu rzeczywistego, aktualizując graf przy każdej nowej interakcji. Identyfikacja podejrzanych połączeń wymaga także umiejętnego wykorzystania agregacji i grupowania wyników według kryteriów domenowych – funkcje ‘GROUP BY’ w Cypher czy ‘.group()’ w Gremlin pozwalają tworzyć listy par obiektów powiązanych przez określony typ relacji wraz z liczbą zaobserwowanych interakcji. Porównanie tych wartości ze średnimi dla całego zbioru ujawnia anomalie statystyczne – nadmierna intensywność kontaktu między dwoma adresami e-mail może wskazywać na nieautoryzowaną transmisję danych.

W scenariuszach wymagających integracji heterogenicznych źródeł danych istotną rolę odgrywa model grafu domenowego łączący trójki RDF z grafem właściwości. Predykaty RDF odwzorowane na typy relacji umożliwiają utrzymanie spójności semantycznej przy importowaniu faktów o kontaktach czy transakcjach z systemów zgodnych ze standardem SPARQL. Pozwala to tworzyć jednolite zapytania mieszane – część danych przeszukiwana jest wzorcami Cyphera, część filtrowana regułami RDF, a wyniki scalane do wspólnego widoku analizowanego pod kątem ryzyka. Warstwa operacyjna identyfikacji często korzysta również z metadanych środowiskowych takich jak historia logowań czy lista urządzeń użytych przy wykonywaniu danej czynności. Połączenie tych informacji z rdzeniem grafu umożliwia np. wykrycie, że wiele kont zalogowało się z tego samego terminala tuż przed serią transferów – taki wzorec można później wykorzystać jako regułę przewidującą wysokie prawdopodobieństwo nadużycia.

Efektywną praktyką jest także użycie filtrów negatywnych eliminujących typowe, nieszkodliwe interakcje (‘WHERE NOT’ w Cypher), co pozwala skupić zasoby obliczeniowe na przypadkach wyjątkowych. W połączeniu z odpowiednio zaprojektowanymi indeksami dla właściwości kluczowych system może identyfikować



podejrzane więzi niemal natychmiast po ich pojawieniu się w bazie. Całościowe podejście do problemu obejmuje więc precyzyjne modelowanie relacji jako pełnoprawnych obiektów wraz ze szczegółowym kontekstem, dynamiczne wzorce zapytań przeszukujące graf pod kątem wielostopniowych powiązań oraz integrację dodatkowych źródeł metadanych umożliwiających korelowanie zachowań technicznych i biznesowych uczestników sieci. Umiejętne łączenie tych mechanizmów przekłada się na zdolność systemu do szybkiego, dokładnego i adaptacyjnego wychwytywania niebezpiecznych powiązań nawet w bardzo rozbudowanych strukturach relacyjnych.

8 Integracja baz grafowych z innymi technologiami

8.1 Integracja z relacyjnymi bazami danych

Integracja grafowych baz danych z relacyjnymi systemami zarządzania bazami danych (RDBMS) wymaga przede wszystkim uwzględnienia zasad spójności strukturalnej i semantycznej pomiędzy odmiennymi modelami przechowywania informacji. Relacyjne bazy danych od dekad stanowią standard w wielu organizacjach, a ich dane są zorganizowane w postaci tabel, gdzie połączenia między rekordami odwzorowuje się poprzez klucze obce oraz mechanizmy integralności referencyjnej (Hernandez). W kontekście integracji z systemem grafowym konieczne jest przekształcenie tych powiązań w krawędzie, przy zachowaniu istotnych właściwości i parametryzacji relacji. Proces translacji zaczyna się często od analizy schematu encja–relacja w relacyjnej bazie i mapowania encji na węzły grafu, zaś relacje, na etykietowane krawędzie. Atrybuty encji stają się właściwościami węzłów, a atrybuty relacji, właściwościami przypisanymi krawędziom. Kardynalność połączeń z diagramów E–R powinna być odwzorowana jako odpowiedni model liczby dopuszczalnych relacji między typami węzłów w grafie.

Takie dopasowanie pozwala zachować logikę biznesową zapisaną wcześniej w mechanizmach integralności, np. ograniczenie do jednej relacji typu OWNS między partnerem a kontraktem można realizować poprzez unikalne indeksy lub reguły walidujące w narzędziu grafowym. Bardzo pomocnym standardem przy integracji jest R2RML (*RDB to RDF Mapping Language*), który umożliwia automatyczne odwzorowanie tabel na zapis trójek RDF. W sytuacjach, gdy docelowa baza grafowa obsługuje model RDF lub hybrydowy graf domenowy, mapowanie takie przenosi dane relacyjne do formatu zgodnego ze SPARQL bez konieczności ręcznego tworzenia krawędzi i węzłów. R2RML pozwala również definiować reguły nazw predykatów – dzięki czemu można utrzymać spójność nazewnictwa między światem



relacyjnym a grafowym. W implementacjach typowego grafu właściwości proces integracyjny przebiega nieco inaczej, odpowiednikiem R2RML stają się dedykowane narzędzia transformujące dane eksportowane z RDBMS (CSV, JSON) w instrukcje CREATE lub MERGE dla silnika grafowego. Niektóre platformy oferują komponenty tzw. ingestion services, które potrafią łączyć się bezpośrednio do SQL-owego źródła, wykonać selekcję danych oraz utworzyć odpowiednie elementy grafu na podstawie ustalonego wzorca.

Dzięki natywnej obsłudze plików wsadowych (jak DataStax Bulk Loader) możliwe jest szybkie masowe przenoszenie struktur tabelarycznych do środowiska grafowego bez znacznego obciążenia klastra produkcyjnego. Integracja wymaga decyzji co do kierunku przepływu danych: czy będzie to migracja jednorazowa ze „starej” bazy relacyjnej do nowej grafowej, czy też systemy będą działały równolegle z synchronizacją wybranych zbiorów rekordów. W przypadku tego drugiego scenariusza należy wdrożyć mechanizmy monitorowania zmian w RDBMS i odwzorowywania ich jako odpowiednich modyfikacji w grafie.

Z technicznego punktu widzenia może to oznaczać implementację procedur wyzwalających (triggers) po stronie SQL, które publikują zdarzenia do kolejki pośredniej, stamtąd trafiają one do procesora integracyjnego aktualizującego strukturę grafową. Należy pamiętać o różnicach semantycznych pomiędzy sposobem obsługi NULL-ów czy braków wartości: podczas gdy w modelu relacyjnym brak wartości może oznaczać nieistnienie więzi lub brak danych o obiekcie, w modelu grafowym nieistnienie krawędzi jest jednoznaczne z brakiem relacji. To powoduje potrzebę dodatkowych reguł interpretacyjnych przy translacji. Integracja danych może obejmować również konwersję indeksów, choć graf właściwości wykorzystuje mechanizm sąsiedztwa bez indeksowego do przechodzenia, to dla operacji punktowych importer może generować indeksy dla właściwości takich jak identyfikatory klientów czy numery zamówień. Taka praktyka pozwala utrzymać porównywalną szybkość wyszukania konkretnego rekordu jak w środowisku SQL.

W przypadku współpracy grafu RDF z relacyjną bazą często korzysta się z narzędzi dostosowujących zapytania – przykładowo SPARQL może być tłumaczony na SQL wybierający dane źródłowe prosto z tabel oraz przekształcający je do formatu trójek. Odwrotna ścieżka polega na wykonywaniu zapytań SQL nad warstwą składowania trójek leżącego „na” silniku RDBMS (np. moduł RDF firmy Oracle). Pozwala to pozostawić dotychczasową infrastrukturę bazodanową i jedynie rozszerzyć jej możliwości o natywne przechowywanie semantycznych relacji.

Scenariusze biznesowe integrujące oba modele obejmują np. raportowanie: dane transakcyjne mogą pozostać zapisane w formacie tabelarycznym dla



kompatybilności z systemami księgowymi, natomiast sieci powiązań klientów i produktów są równolegle budowane jako graf służący analizie rekomendacyjnej albo detekcji nadużyć. W takim układzie ważne jest utrzymanie jednolitych kluczy głównych / identyfikatorów elementów pomiędzy obiema bazami, aby możliwe było dwukierunkowe mapowanie wyników analiz. Dobrą praktyką jest też prowadzenie warstwy abstrakcji API zapewniającej jednolite metody dostępu niezależnie od źródła – zapytanie o historię kontaktów klienta może korzystać zarówno ze schematu SQL jak i przechodzenia Cyphera czy Gremlina. Warstwa ta izoluje aplikację od szczegółów implementacyjnych obu modeli oraz ułatwia przyszłe zmiany technologii po jednej ze stron integracji. Współdziałanie RDBMS i bazy grafowej zwiększa elastyczność organizacji: pozwala wykorzystywać mocne strony obu technologii, stabilność i wydajność operacyjną modelu tabelarycznego tam, gdzie potrzebne są transakcje masowe oraz ekspresyjność modeli powiązań oferowaną przez strukturę grafową przy analizach wielopoziomowych zależności. Umiejętne zaprojektowanie procesu wymiany informacji oraz konwersji modeli gwarantuje spójny obraz danych niezależnie od tego, czy analiza odbywa się po stronie SQL czy struktury grafikowej.

8.2 Integracja z systemami NoSQL

Integracja baz grafowych z systemami NoSQL wynika z faktu, że obie technologie należą do szerokiej kategorii magazynów nierelacyjnych, lecz różnią się strukturą danych i optymalizacjami pod kątem typowych operacji. NoSQL obejmuje modele takie jak dokumentowe, klucz–wartość, kolumnowe czy grafowe. W praktyce oznacza to możliwość współdzielenia infrastruktury lub ścieżek przetwarzania danych pomiędzy systemami o odmiennym podejściu do reprezentacji relacji. W kontekście integracji istotne jest rozpoznanie punktów styku – w modelach dokumentowych relacja między rekordami bywa odwzorowana przez identyfikatory przechowywane wewnątrz dokumentu; w środowisku grafowym te same powiązania stają się krawędziami z własnymi właściwościami kontekstowymi, co daje większą ekspresyjność.

Praktycznym rozwiązaniem są hybrydowe architektury, w których fragment danych pozostaje w oryginalnym systemie NoSQL (np. w Cassandra czy MongoDB), a część jest replikowana lub odwzorowywana do bazy grafowej w celu realizacji zadań eksploracyjnych. Integracja może odbywać się poprzez regularną synchronizację wsadów – eksport danych z tabel kolumnowych lub struktur JSON i ich translacja do formatu grafu właściwości – albo poprzez mechanizmy strumieniowe reagujące na zdarzenia po stronie źródła. W tym drugim przypadku stosuje się kolejki



komunikatów lub systemy typu Kafka, które przechwytyją zapis/aktualizację w magazynie NoSQL i przekazują ją do procesora aktualizującego strukturę grafową.

Ważnym aspektem jest utrzymanie spójności semantycznej przy mapowaniu typów danych. Modele grafowe wymagają jawnej definicji etykiet węzłów i relacji; dane dokumentowe mogą nie zawierać tak wyodrębnionej klasyfikacji. Podczas integracji konieczne jest więc nadanie dokumentom ról odpowiadających etykietom – np. User, Product, Transaction. Takie grupowanie umożliwia późniejsze tworzenie indeksów lokalnych dla właściwości kluczowych i precyzyjne filtrowanie wyników. Z kolei modele kolumnowe oparte na szerokich tabelach pozwalają odwzorować wiele właściwości elementu grafowego jako zestaw kolumn – co upraszcza migrację przy dobrze zdefiniowanym schemacie.

Z perspektywy mechaniki przechowywania, istotne jest zsynchronizowanie strategii fragmentowania między systemem NoSQL a bazą grafową. NoSQL horyzontalnie dzieli dane według klucza partycji; dla grafu wyzwaniem jest minimalizacja liczby krawędzi przekraczających granice partycji. W integracjach dąży się do wspólnego planu podziału zbioru tak, aby powiązane węzły trafiły na ten sam serwer również po stronie NoSQL – eliminuje to opóźnienia wynikające z sieciowej rekonstrukcji ścieżek. Problem dopasowania języków zapytań ma wymiar techniczny i użytkowy. Aplikacja pracująca dotychczas z interfejsem agregatowym (query API znane z dokumentowych lub kolumnowych baz) musi zaadaptować się do wzorcowego języka grafowego (np. Cypher czy Gremlin). Możliwe jest używanie tłumaczy zapytań – część konstrukcji filtrujących z zapytań MongoDB może zostać przełożona na 'MATCH' + 'WHERE' w Cypher przy zachowaniu logiki warunków.

Podobnie operacje agregujące ('groupBy') można odwzorować na '.group()' Gremlina. W hybrydach pojawia się też opcja delegowania fragmentu zapytania do silnika NoSQL celem redukcji zbioru wejściowego przed przechodzeniem po stronie grafowej. Integracje obejmują często rozszerzenia modelu o metadane β-
posługiwanie mechanizmami krawędzi jako węzeł dla obiektów specyficznych tylko dla jednej ze stron współpracy. Przykładem mogą być relacje opisujące sesję użytkownika stworzone przez system aplikacyjny zapisujący logi w Cassandra – te zdarzenia trafiają jako osobne węzły powiązane krawędziami do konta użytkownika po stronie grafu. Dzięki temu można analizować historię wraz z danymi transakcyjnymi bez przenoszenia całości logów.

Szczególnie ciekawa jest integracja grafu właściwości z trójką RDF obsługiwanym przez silnik działający w warstwie NoSQL. Grafy domenowe pozwalają utrzymać interoperacyjność semantyczną między predykatami RDF a typami krawędzi grafowych, co sprzyja scalaniu różnych repozytoriów danych bez utraty



kompatybilności ze SPARQL oraz Cypher. Taka dwutorowość dostępu umożliwia kontynuowanie analiz ontologicznych przy jednoczesnym korzystaniu z wydajnych przechodzeń. Warto uwzględnić także rolę dedykowanych narzędzi integracyjnych oferujących wsparcie dla popularnych ekosystemów – przykładowo Neo4j może łączyć się bezpośrednio z MongoDB poprzez konektory, które synchronizują zmiany kolekcji jako aktualizacje subgrafów odpowiadających tym dokumentom. Również platformy typu multi-model (OrientDB) pozwoliły historycznie na natywne obsłużenie zarówno trybu dokumentowego, jak i grafowego we wspólnym magazynie, co upraszcza integrację logiczną, choć nie zawsze zapewnia pełną specjalizację wydajnościową każdego modelu. Istotna jest kontrola kosztów operacyjnych, częste aktualizacje danych wyjściowych mogą obciążać oba systemy podczas synchronizacji. Dlatego zwykle stosuje się mechanizmy różnicowe (delta sync), które przesyłają jedynie zmienione fragmenty zamiast pełnego zbioru. Po stronie grafowej ułatwia to utrzymanie spójności struktury i właściwości bez ponownego ładowania dużych partii danych. Integracja baz grafowych i systemów NoSQL łączy zalety obu technologii: dynamikę luźno ustrukturyzowanych rekordów oraz ekspresję modeli relacyjnych bogatych topologicznie. Umiejętne wykorzystanie wspólnych cech (rozproszone przechowywanie, elastyczne schematy) oraz redukcja różnic poprzez odpowiednie mapowanie struktur zwiększa efektywność całego rozwiązania i pozwala budować kompleksowe środowiska analityczne pracujące równolegle na różnych typach modeli danych zgodnie ze specyfiką poszczególnych źródeł.

9 Przyszłość baz grafowych

9.1 Nowe kierunki badań

Kierunki przyszłych badań nad bazami grafowymi wydają się coraz wyraźniej zmierzać ku pogłębionej integracji z koncepcją semantycznej analizy danych oraz automatyzacji procesu odkrywania wiedzy w strukturach powiązań. W środowiskach produkcyjnych dostrzegalny jest ruch w stronę tzw. wykresów wiedzy o zdarzeniach opartych na danych, które łączą bieżące strumienie zdarzeń z uporządkowanym modelem grafowym. Pozwala to nie tylko śledzić dynamikę relacji pomiędzy bytami, ale też osadzać je w szerszym kontekście biznesowym czy technologicznym, w którym zdarzenia mają określone następstwa i zależności. Badania nad tego typu systemami wymagają dopracowania mechanizmów synchronizacji strumieni ze stanem bazowym grafu – musi on pozostawać spójny pod względem topologii i metadanych mimo napływu tysięcy nowych krawędzi na sekundę.



Rozwój architektur hybrydowych łączących model grafu właściwości z RDF i grafami domenowymi tworzy pole do eksperymentów nad współdzielonym rdzeniem semantycznym. Tego rodzaju badania skupiają się na tym, jak wykorzystać zalety bogatych atrybutów grafu właściwości przy jednoczesnej zgodności ze standardami SPARQL oraz OWL. Możliwość traktowania krawędzi jako pełnoprawnych węzłów (krawędź jako węzeł) otwiera nowe scenariusze modelowania relacji drugiego rzędu, takich jak opis kontraktu zawierającego metadane dotyczące negocjacji, powiązań osobowych uczestników czy zastosowanych klauzul prawnych. Badania w tym kierunku wymagają jednak udoskonalenia procesów translacji zapytań między różnymi składniami, aby nie tracić wydajności przechodzenia przy zachowaniu kompatybilności semantycznej.

Kolejnym obszarem intensywnej eksploracji jest automatyczna inferencja logiczna w grafach wiedzy. Wykorzystanie reguł generujących nowe fakty na podstawie istniejących relacji zwiększa aktualność i użyteczność baz bez konieczności ręcznych aktualizacji. W badaniach koncentruje się na optymalizacji kosztu takich operacji, szczególnie gdy działają one w czasie rzeczywistym i muszą współdziałać z transakcyjnym rdzeniem bazy spełniającym wymagania ACID. Istotne są także zagadnienia związane z przechowywaniem wyników inferencji – muszą one być spójne ze strukturą bazową, a jednocześnie łatwo identyfikowalne jako pochodne istniejących faktów. Wraz ze wzrostem skali danych rośnie zapotrzebowanie na nowe techniki przetwarzania rozproszonego grafów. Problem partycjonowania przy minimalizacji liczby przecinanych krawędzi pozostaje wyzwaniem badawczym, zwłaszcza gdy struktura połączeń ulega dynamicznym zmianom.

Obecne algorytmy fragmentowania rzadko radzą sobie efektywnie z przypadkami silnie skorelowanymi, takimi jak sieci społecznościowe czy systemy rekomendacyjne, gdzie lokalna gęstość połączeń bywa bardzo nierównomierna. Z tego powodu badania kierują się ku adaptacyjnym metodom rozmieszczania węzłów opartym na analizie bieżących wzorców dostępu – tak, by minimalizować koszt przechodzenia najczęściej wykonywanych przez aplikacje. Wydajne algorytmy wyszukiwania ścieżek o zmiennej długości oraz regularne zapytania o ścieżkę stają się kolejnym strategicznym celem prac rozwojowych. Potrzeba ta wynika m.in. z zastosowań analityki sieciowej czy wykrywania społeczności, gdzie trawersy obejmują nieprzewidzianą wcześniej liczbę skoków lub wymagają analizy meta faktów, relacji opisujących inne relacje. Wzmocnienie obsługi takich scenariuszy powinno objąć zarówno warstwę językową (rozszerzenie składni Cyphera czy Gremlina), jak i warstwę planowania zapytań zdolną zoptymalizować wykonanie wieloetapowych przebiegów po grafie.



Badania nad integracją baz grafowych z uczeniem maszynowym są coraz częściej związane z tworzeniem modeli Graph Neural Networks (GNN), które adaptują reprezentacje osadzania wierzchołków i krawędzi do przewidywania nowych połączeń lub klasyfikowania elementów sieci. Grafy wiedzy z ontologiami szczególnie dobrze nadają się do trenowania GNN dzięki bogatemu kontekstowi semantycznemu, jednak wymaga to opracowania metod skalowania przetwarzania tensora sąsiedztwa dla bardzo dużych zbiorów danych bez utraty jakości predykcji. W tym obszarze istotną rolę odgrywa badanie mechanizmów próbkowania podgrafów oraz inkrementalnego uczenia przy stale rosnącej bazie. Warto również wskazać rozwój języków zapytań hybrydowych zdolnych płynnie łączyć paradygmat deklaratywny i proceduralny. Prace zmierzają do stworzenia uniwersalnej składni umożliwiającej zarówno wzorce typu MATCH z filtrami właściwości, jak i kroki przechodzenia sterowane logiką warunkową znaną z Gremlina. Realizacja takiej wizji mogłaby ułatwić migrację kodu analitycznego pomiędzy silnikami oraz wspierać zespoły korzystające równocześnie z wielu środowisk grafowych. Zarysowane kierunki badań obejmują również zagadnienia interoperacyjności między źródłami heterogenicznymi, konieczne jest doskonalenie narzędzi mapujących struktury danych o różnym stopniu formalizacji (od dokumentowych NoSQL po silnie typowane RDBMS) na spójny model grafowy bez utraty metadanych krytycznych dla analizy. Takie rozwiązania mogą korzystać ze standardów jak R2RML dla mapowań do RDF oraz autorskich schematów translacji do grafu właściwości. Dalsze prace przewidują dodanie warstwy weryfikacyjnej sprawdzającej integralność semantyczną importowanych danych przed ich osadzeniem w rdzeniu grafu. Podsumowując, przyszłe prace skoncentrują się na synergii między natywnym przechowywaniem i przetwarzaniem grafu a technologiami semantycznymi oraz uczeniem maszynowym; na poprawie skalowalności poprzez inteligentne partycjonowanie; rozszerzeniu ekspresji języków zapytań o bardziej elastyczne schematy przechodzenia; oraz wzmocnieniu zdolności integracyjnych baz grafowych jako centralnego repozytorium wiedzy organizacyjnej przystosowanego do pracy ze strumieniami wydarzeń i danymi wieloformatowymi.

9.2 Wpływ sztucznej inteligencji na rozwój baz grafowych

Rzeczywisty rozwój baz grafowych coraz częściej jest sprzężony z postępem w dziedzinie sztucznej inteligencji, a zwłaszcza uczenia maszynowego. Relacje między tymi obszarami można rozpatrywać w dwóch kierunkach: AI jako odbiorca i konsument danych grafowych oraz AI jako czynnik kształtujący technologię i architekturę samych baz. W roli odbiorcy, algorytmy uczące się wykorzystują grafy wiedzy do pozyskiwania kontekstu semantycznego, integracji niejednorodnych źródeł informacji



i wzbogacania modeli predykcyjnych o strukturalne zależności. Tego typu grafy mogą pełnić funkcję „mózgu” organizacji, centralnego repozytorium, w którym powstaje spójny obraz bytów biznesowych oraz powiązań między nimi. Dzięki temu nawet niewielkie porcje danych treningowych stają się wartościowe dla modelu, gdy są osadzone w szerszej sieci relacji umożliwiającej wyprowadzenie dodatkowych cech opisowych na bazie ścieżek połączeń.

Z drugiej strony sztuczna inteligencja bywa czynnikiem modyfikującym technologię przechowywania i przetwarzania grafów. Jednym z kluczowych trendów jest automatyzacja etapów przygotowania danych, w infrastrukturze DataOps algorytmy ML wspierają proces „groomingu” zbiorów przez semantyczne wzbogacanie, oczyszczanie i standaryzację. Automatyczne procedury pozwalają na bieżąco integrować dane z wielu systemów, a następnie odwzorowywać je jako spójny graf wraz z metadanymi koniecznymi dla późniejszych analiz AI. Zastosowanie takich mechanizmów skraca czas od momentu dodania nowej informacji w źródle do chwili jej pełnej dostępności dla modeli predykcyjnych, co jest istotne w środowiskach wymagających natychmiastowej reakcji.

AI wpływa też na metody eksploracji danych grafowych, rozwijane są narzędzia łączące zapytania strukturalne z analizą treści nieustrukturyzowanej, np. wyszukiwanie semantyczne dokumentów i wiadomości powiązanych z określonym bytem lub grupą. Systemy te potrafią rozpoznać synonimy, konteksty branżowe czy ukryte korelacje i powiązać je z odpowiednimi częściami grafu, co rozszerza zakres możliwych zapytań i umożliwia wykrywanie zależności wcześniej trudnych do uchwycenia klasycznymi metodami. Istotnym efektem synergii AI–graf jest rozwój modeli typu Graph Neural Networks (GNN), które pozwalają bezpośrednio uczyć się reprezentacji węzłów i krawędzi na podstawie ich topologicznego otoczenia oraz właściwości.

Bazy grafowe muszą dostosować swoje API i sposoby przechowywania tak, aby zapewnić szybki dostęp do lokalnych sąsiedztw i sprawną serializację partii danych przekazywanych do trenowania sieci neuronowej. W tym kontekście architektura bazy wpływa na jakość i efektywność procesu uczenia, mechanizm sąsiedztwa bez indeksowego staje się fundamentem dla wydajnego zasysania podgrafów potrzebnych przy losowym próbkowaniu czy przetwarzaniu mini-partii. Sztuczna inteligencja modyfikuje również logikę inferencji wewnątrz samych systemów grafowych. Reguły generowania nowych faktów oparte na modelach predykcyjnych zastępują bądź uzupełniają tradycyjne systemy reguł logicznych. Na przykład klasyczny mechanizm dodający relacje na podstawie prostych wzorców może zostać rozszerzony o element oceny prawdopodobieństwa, wyliczany przez model ML. Dzięki temu graf rozwija się „inteligentnie”, tworząc połączenia tylko wtedy, gdy ich



istnienie jest potwierdzone zarówno przez formalne reguły, jak i predykcję wynikającą z danych historycznych.

Integracja AI wymusza też zmiany w językach zapytań oraz narzędziach analitycznych baz grafowych. Wprowadzane są konstrukcje umożliwiające uruchamianie modeli bezpośrednio wewnątrz cyklu zapytania czy aktualizacji danych, przykładowo wywołanie modelu klasyfikacyjnego na każdym nowo dodanym węźle klienta może automatycznie przypisać mu kategorię ryzyka lub potencjału zakupowego. Takie podejście eliminuje konieczność oddzielnej fazy analizy poza bazą, skracając ścieżkę od pojawienia się danych do uzyskania wartości biznesowej. Automatyczne rekomendacje dotyczą także projektowania samego schematu grafu pod kątem przewidywanego użycia przez AI. Algorytmy mogą sugerować agregację pewnych rodzajów relacji lub dodanie brakujących typów krawędzi na podstawie analizy dostępnych zapytań oraz ich skuteczności modelowej. Przykładem jest automatyczne wykrycie potrzeby dodania warstwy meta-relacji opisujących dynamikę interakcji użytkownika z produktem (częstotliwość zmian preferencji), jeśli modele rekomendacyjne wskazują dużą wagę tej cechy dla poprawności przewidywań.

W kontekście architektury systemowej integracja graf–AI obejmuje często strumieniowe ładowanie danych ze źródeł transakcyjnych do rdzenia bazy wraz z natychmiastową aktualizacją modelu ML (Author, b). Wymaga to mechanizmów gwarantujących spójność pomiędzy stanem bieżącym a używanym przez model zestawem cech, spójność ta musi być zachowana zarówno przy pracy w trybie czasu rzeczywistego jak i seryjnym. Wpływ sztucznej inteligencji widoczny jest również w rosnącym znaczeniu semantycznego zarządzania danymi: AI korzysta z ontologii zapisanych w grafie wiedzy do ulepszania interpretacji wyników oraz generowania bardziej trafnych odpowiedzi przy interfejsach typu chatbot. Im bogatszy słownik pojęć i definicji skojarzonych z elementami grafu, tym większe możliwości adaptacyjne ma aplikacja AI wobec zmieniającego się kontekstu pytań użytkownika.

Całość przemian prowadzi do sytuacji, gdzie granica pomiędzy silnikiem bazodanowym a platformą przetwarzania AI staje się mniej wyraźna: bazy grafowe zaczynają działać jako aktywny komponent ekosystemu analitycznego, nie tylko przechowując dane, ale też uczestnicząc w ich interpretacji poprzez wywołania modeli uczenia maszynowego oraz własną logikę inferencyjną wspartą AI. Taki kierunek rozwoju zwiększa rolę technologii graficznych jako centralnej warstwy integracyjno-analitycznej nowoczesnych organizacji.



10 Zakończenie

Analiza baz grafowych ukazuje ich fundamentalną rolę w nowoczesnym przetwarzaniu danych, zwłaszcza tam, gdzie kluczowe są relacje i powiązania między obiektami. Model grafowy, oparty na wierzchołkach i krawędziach, umożliwia naturalne odwzorowanie złożonych struktur, które trudno efektywnie reprezentować w tradycyjnych systemach relacyjnych. Mechanizm sąsiedztwa bez indeksowego oraz natywna architektura magazynów grafowych zapewniają wysoką wydajność operacji eksploracyjnych, co jest szczególnie istotne przy analizie gęsto powiązanych danych, takich jak sieci społeczne, systemy rekomendacyjne czy sieci powiązań biznesowych.

Wprowadzenie języków zapytań takich jak Cypher i Gremlin umożliwia intuicyjne formułowanie zarówno prostych, jak i zaawansowanych wzorców eksploracji grafu, łącząc deklaratywność z elastycznością proceduralną. Rozwój hybrydowych modeli, łączących graf właściwości z RDF i grafami domenowymi, pozwala na integrację różnorodnych źródeł danych oraz zachowanie zgodności ze standardami semantycznymi, co jest kluczowe dla interoperacyjności i rozbudowanych analiz wiedzy. Wydajność systemów grafowych zależy od świadomego projektowania schematów, optymalizacji zapytań oraz efektywnego indeksowania, które razem minimalizują koszty operacyjne i skracają czas odpowiedzi.

Zastosowania baz grafowych obejmują szeroki zakres dziedzin, od analizy połączeń i wykrywania społeczności w sieciach społecznościowych, przez systemy rekomendacyjne, aż po analizę sieci powiązań biznesowych oraz wykrywanie oszustw. W każdym z tych obszarów model grafowy pozwala na uchwycenie i analizę relacji w sposób bardziej naturalny i efektywny niż tradycyjne podejścia, co przekłada się na lepszą jakość wyników i szybsze reakcje systemów. Integracja z relacyjnymi bazami danych oraz systemami NoSQL umożliwia wykorzystanie istniejących zasobów i technologii, jednocześnie rozszerzając możliwości analityczne o funkcje charakterystyczne dla grafów.

Perspektywy rozwoju baz grafowych wiążą się z pogłębioną integracją z technologiami sztucznej inteligencji, w tym uczeniem maszynowym i sieciami neuronowymi działającymi na grafach. Automatyzacja procesów przygotowania danych, wzbogacanie semantyczne oraz dynamiczna inferencja faktów stanowią kierunki, które mogą znacznie zwiększyć efektywność i zakres zastosowań tych systemów. Wyzwania związane ze skalowalnością, adaptacyjnym partycjonowaniem oraz rozszerzeniem ekspresji języków zapytań wymagają dalszych badań i innowacji, które pozwolą sprostać rosnącym wymaganiom środowisk produkcyjnych.



Fundusze Europejskie
dla Rozwoju Społecznego



Rzeczpospolita
Polska

Dofinansowane przez
Unię Europejską



W efekcie bazy grafowe stają się integralnym elementem nowoczesnych ekosystemów danych, łącząc elastyczność modelowania, wydajność przetwarzania oraz zdolność do integracji heterogenicznych źródeł informacji. Ich rozwój i adaptacja do potrzeb analityki, zarządzania wiedzą oraz systemów wspomagających decyzje przyczyniają się do tworzenia bardziej inteligentnych i responsywnych rozwiązań informatycznych, które odpowiadają na wyzwania współczesnych organizacji i społeczeństw.





11. Literatura

1. Hyman J. A., Massaron L., McFedries P., (2024) *Data analytics & visualization all-in-one for dummies*, Wyd. For Dummies.
2. *The Big Book of Data Engineering*. (2024), Databricks Inc.
3. Blumauer A., Nagy H., (2020) *The knowledge graph cookbook*. Wyd. edition mono/monochrom
4. Hřivnáč J., (2020), *Using graph databases*, EPJ Web of Conferences. Available at: <https://doi.org/10.1051/epjconf/202024504004>.
5. Gosnell D.K., Broecheler M. (2021), *Dane grafowe w praktyce: Jak technologie grafowe ułatwiają rozwiązywanie złożonych problemów*, Wyd. Helion
6. Hernandez, M.J. (2014), *Projektowanie baz danych dla każdego: Przewodnik krok po kroku*. Wyd. Helion
7. Patil, S., Vaswani, G. and Bhatia, A. (2014) *Graph databases- an overview*, International Journal of Computer Science and Information Technologies. <https://www.ijcsit.com/docs/Volume%205/vol5issue01/ijcsit20140501141.pdf>
8. Pokorný, J. (2015) *Graph databases: Their power and limitations*, in. Available at: https://doi.org/10.1007/978-3-319-24369-6_5.
9. Ramachandran, S. (2015) *Graph Database Theory*. Wyd. LambdaZen <https://www.lambdazen.com/assets/pdf/GraphDatabaseTheory.pdf>
10. Rodrigues C., Khanchandani A. (2023). *Performance Comparison of Graph Database and Relational Database*. 10.13140/RG.2.2.27380.32641.
11. Robinson I., Webber J., Eifrem E. (2015) *Graph Databases: New Opportunities for Connected Data (2nd. ed.)*. O'Reilly Media, Inc.
12. Sasaki, B.M., Chao, J. and Howard, R. (2018) *Graph databases for beginners*. Wyd. Neo4j
13. Tittel E. (2021) *Data fabric*, Hitachi Vantara Special Edition.
14. Vrgoc, D. et al. (2023) 'MillenniumDB: An open-source graph database system', Data Intelligence, 5(3), pp. 560–610. Available at: https://doi.org/10.1162/dint_a_00209.

