

Snowflake

**Nowoczesna
inżynieria danych
w praktyce**

Maja Ferle



Tytuł oryginału: Snowflake Data Engineering

Tłumaczenie: Ksawery Sosnowski

ISBN: 978-83-289-2898-5

© Helion S.A. 2026.

Authorized translation of the English edition © 2025 Manning Publications.

This translation is published and sold by permission of Manning Publications, the owner of all rights to publish and sell the same.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from the Publisher.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiejkolwiek postaci jest zabronione. Wykonywanie kopii metodą kserograficzną, fotograficzną, a także kopiowanie książki na nośniku filmowym, magnetycznym lub innym powoduje naruszenie praw autorskich niniejszej publikacji.

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autor oraz wydawca dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autor oraz wydawca nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Drogi Czytelniku!

Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres

helion.pl/user/opinie/snowfl

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Helion S.A.

ul. Kościuszki 1c, 44-100 Gliwice

tel. 32 230 98 63

e-mail: helion@helion.pl

WWW: helion.pl (księgarnia internetowa, katalog książek)

Printed in Poland.

- [Kup książkę](#)
- [Poleć książkę](#)
- [Oceń książkę](#)

- [Księgarnia internetowa](#)
- [Lubię to! » Nasza społeczność](#)

Spis treści

<i>Przedmowa</i>	11
<i>Wstęp</i>	13
<i>Podziękowania</i>	15
<i>O książce</i>	17
<i>Informacje o autorce</i>	22
CZĘŚĆ I. WPROWADZENIE DO INŻYNIERII DANYCH Z WYKORZYSTANIEM PLATFORMY SNOWFLAKE	23
1. <i>Inżynieria danych w środowisku Snowflake</i>	25
1.1. Snowflake a inżynieria danych	26
1.1.1. <i>Architektura Snowflake’a</i>	26
1.1.2. <i>Funkcje Snowflake’a używane w inżynierii danych</i>	28
1.2. Obowiązki inżyniera danych w Snowflake’u	29
1.2.1. <i>Pozyskiwanie danych z systemów źródłowych</i>	31
1.2.2. <i>Przekształcanie danych</i>	32
1.2.3. <i>Udostępnianie danych odbiorcom końcowym</i>	33
1.2.4. <i>Stosowanie podstawowych składników</i>	33
1.3. Budowanie potoków danych	37
1.4. Inżynieria danych w zastosowaniach Snowflake’a	38
Podsumowanie	40
2. <i>Tworzenie pierwszego potoku danych</i>	42
2.1. Konfigurowanie konta Snowflake’a	44
2.2. Umieszczanie pliku CSV w etapie	45

2.3. Wczytywanie danych z pliku przejściowego do tabeli docelowej ...	48
2.3.1. Wczytywanie danych z pliku przejściowego do tabeli przejściowej	49
2.3.2. Scalanie danych z tabeli przejściowej w tabelę docelową	51
2.4. Przekształcanie danych za pomocą poleceń SQL	53
2.5. Automatyzowanie procesu za pomocą zadań	55
Podsumowanie	59

CZĘŚĆ II. WPROWADZANIE, PRZEKSZTAŁCANIE I PRZECHOWYWANIE DANYCH 61

3. <i>Dobre praktyki przygotowywania danych</i>	63
3.1. Tworzenie etapów zewnętrznych	66
3.1.1. Konfigurowanie integracji magazynu danych	67
3.1.2. Tworzenie etapu zewnętrznego z użyciem integracji magazynu danych	69
3.1.3. Tworzenie etapu zewnętrznego z użyciem poświadczeń	70
3.1.4. Wczytywanie danych z plików przejściowych do tabeli przejściowej	71
3.1.5. Zapobieganie powielaniu danych podczas wczytywania z plików przejściowych	74
3.1.6. Używanie nazwanego formatu pliku	75
3.2. Przeglądanie metadanych etapów przy użyciu tabel katalogowych	77
3.3. Przygotowywanie plików danych do efektywnego wprowadzania	78
3.3.1. Zalecenia dotyczące rozmiaru plików	78
3.3.2. Organizowanie danych według ścieżek	79
3.4. Tworzenie potoków danych z wykorzystaniem tabel zewnętrznych	80
3.4.1. Kwerendowanie danych z etapów zewnętrznych za pomocą tabel zewnętrznych	81
3.4.2. Wykorzystywanie widoków zmaterializowanych do poprawy wydajności kwerend	83
Podsumowanie	84
4. <i>Przekształcanie danych</i>	86
4.1. Pozyskiwanie częściowo ustrukturyzowanych danych z magazynu chmurowego	88
4.1.1. Tworzenie integracji magazynu danych	89
4.1.2. Tworzenie etapu zewnętrznego	91

4.1.3. Analiza struktury JSON	91
4.1.4. Wczytywanie danych JSON do typu VARIANT	92
4.2. Spłaszczanie danych półstrukturalnych do tabel relacyjnych	94
4.3. Hermetyzacja przekształceń za pomocą procedur składowanych	98
4.3.1. Tworzenie prostej procedury składowanej	101
4.3.2. Uwzględnianie zwracanej wartości w procedurze składowanej	102
4.3.3. Implementacja obsługi wyjątków w procedurach składowanych	103
4.4. Dodawanie rejestrowania zdarzeń do procedur składowanych ...	104
4.5. Budowanie niezawodnych potoków danych	107
Podsumowanie	110
5. Ciągłe wprowadzanie danych	113
5.1. Porównanie masowego i ciągłego pobierania danych	115
5.2. Przygotowywanie plików w magazynie chmurowym	116
5.2.1. Tworzenie integracji magazynu danych	117
5.2.2. Tworzenie etapu zewnętrznego	118
5.3. Konfigurowanie usługi Snowpipe z powiadomieniami w chmurze	120
5.3.1. Konfigurowanie komunikatów Event Grid dla zdarzeń magazynu obiektów blob	121
5.3.2. Tworzenie integracji powiadomień	124
5.3.3. Tworzenie obiektu pipe	125
5.3.4. Ciągłe pobieranie danych	127
5.3.5. Spłaszczanie struktury JSON do formatu relacyjnego	128
5.4. Przekształcanie danych za pomocą tabel dynamicznych	129
Podsumowanie	132
6. Natywne wykonywanie kodu z wykorzystaniem usługi Snowpark	134
6.1. Wprowadzenie do usługi Snowpark	136
6.2. Tworzenie procedury Snowpark w arkuszu	137
6.3. Korzystanie z interfejsu API SQL w lokalnym środowisku programistycznym	142
6.3.1. Instalowanie i konfigurowanie lokalnego środowiska programistycznego	142
6.3.2. Tworzenie sesji Snowflake'a	143
6.3.3. Przechowywanie danych uwierzytelniających w pliku konfiguracyjnym	144
6.3.4. Obsługa kwerend danych i poleceń SQL	145

6.4. Tworzenie wymiaru dat w usłudze Snowpark Python	146
6.5. Praca z ramkami danych	148
6.6. Wczytywanie danych z pliku CSV do tabeli Snowflake'a	151
6.7. Przekształcanie danych przy użyciu ramek danych	155
Podsumowanie	157
7. Rozszerzanie zestawów danych wynikami z dużych modeli językowych	159
7.1. Konfigurowanie dostępu do sieci zewnętrznej	161
7.2. Wywoływanie punktu końcowego API z poziomu usługi Snowpark	164
7.2.1. Tworzenie funkcji UDF do pobierania opinii klientów	165
7.2.2. Interpretacja wyników funkcji UDF	167
7.2.3. Przechowywanie opinii w tabeli	169
7.3. Analiza nastrojów w opiniach klientów	170
7.4. Interpretacja e-maili z zamówieniami przy użyciu modeli językowych w celu zaoszczędzenia czasu	172
7.4.1. Tworzenie procedury składowanej do interpretowania e-maili od klientów	173
7.4.2. Tworzenie zapytania	174
7.4.3. Zapisywanie wyniku w tabeli CSV	175
7.4.4. Ocena wyników	176
Podsumowanie	178
8. Optymalizowanie wydajności kwerend	180
8.1. Pobieranie danych z platformy Snowflake Marketplace	181
8.2. Analiza danych geograficznych	184
8.2.1. Funkcje geograficzne Snowflake'a	184
8.2.2. Kopiowanie danych z udostępnionej bazy danych	186
8.2.3. Przeglądanie parametrów wykonania zapytań przy użyciu profilu kwerendy	188
8.3. Mikropartycje w Snowflake'u	190
8.3.1. Koncepcyjny przykład mikropartycji	190
8.3.2. Selektywne pomijanie mikropartycji	192
8.4. Optymalizacja magazynu pamięci za pomocą klasteryzacji	194
8.4.1. Przeglądanie informacji o klasteryzacji	194
8.4.2. Dodawanie kluczy klasteryzacji do tabeli	195
8.4.3. Monitorowanie procesu klasteryzacji	196
8.4.4. Analiza poprawy wydajności kwerend po klasteryzacji	198

8.5. Zwiększanie wydajności kwerend poprzez optymalizację wyszukiwania	199
8.5.1. Dodawanie optymalizacji wyszukiwania do tabeli	201
8.5.2. Analiza wydajności kwerend po zastosowaniu optymalizacji wyszukiwania	201
8.6. Ogólne wskazówki dotyczące poprawy wydajności kwerend	203
8.6.1. Pisanie wydajnych kwerend SQL	203
8.6.2. Identyfikowanie kwerend nadających się do optymalizacji	204
Podsumowanie	205
9. Kontrolowanie kosztów	207
9.1. Kwestia kosztów w Snowflake'u	208
9.1.1. Całkowity koszt korzystania ze Snowflake'a	209
9.1.2. Koszty zasobów obliczeniowych	210
9.1.3. Kredyty magazynu wirtualnego	211
9.2. Dobór rozmiaru wirtualnych magazynów danych	212
9.2.1. Korzystanie z utrwalonych wyników kwerend	215
9.2.2. Porównywanie statystyk kwerend dla magazynów danych o różnych rozmiarach	216
9.2.3. Optymalizacja wydajności kwerend w celu ograniczenia przelewania danych	218
9.3. Optymalizacja wydajności za pomocą buforowania danych	221
9.3.1. Prezentacja pamięci podręcznej metadanych	222
9.3.2. Efektywne wykorzystanie pamięci podręcznej magazynu danych	222
9.4. Redukowanie kolejkowania zapytań	223
9.4.1. Analiza kolejkowania	224
9.4.2. Ograniczanie liczby jednocześnie wykonywanych kwerend	225
9.5. Monitorowanie zużycia mocy obliczeniowej	226
Podsumowanie	227
10. Zarządzanie danymi i kontrola dostępu	230
10.1. Kontrola dostępu oparta na rolach	231
10.1.1. Gotowe role systemowe	232
10.1.2. Role niestandardowe	232
10.1.3. Projektowanie modelu RBAC	233
10.2. Zabezpieczanie danych za pomocą zasad dostępu do wierszy ...	239
10.3. Ochrona danych za pomocą zasad maskowania	243
Podsumowanie	245

CZĘŚĆ III. BUDOWANIE POTOKÓW**PRZETWARZANIA DANYCH247****11. Projektowanie potoków danych 249**

11.1. Projektowanie potoków danych 250

11.1.1. Wyodrębnianie danych 251

11.1.2. Porównanie wzorców potoków danych 252

11.1.3. Wybór warstw przekształcania danych 254

11.1.4. Organizacja warstw przechowywania danych 256

11.1.5. Tworzenie schematów z kontrolą dostępu 258

11.2. Budowanie przykładowego potoku danych 259

11.2.1. Implementacja warstwy wydobywania 259

11.2.2. Implementacja warstwy przygotowywania etapów 261

11.2.3. Implementacja warstwy hurtowni danych 263

11.2.4. Implementacja warstwy tworzenia raportów 265

Podsumowanie 266

12. Przyrostowe pobieranie danych 268

12.1. Porównanie metod pobierania danych 270

12.1.1. Pełne pobieranie danych 270

12.1.2. Przyrostowe pobieranie danych 271

12.2. Przechowywanie historii z wykorzystaniem
wymiarów wolnozmiennych 271

12.2.1. Wymiary wolnozmiennych typu drugiego 272

12.2.2. Strategia dopisywania 273

12.2.3. Projektowanie idempotentnych potoków danych 275

12.3. Wykrywanie zmian za pomocą strumieni Snowflake'a 275

12.3.1. Przyrostowe pobieranie plików
z magazynu chmurowego 27612.3.2. Przechowywanie historii w trakcie
przyrostowego pobierania danych 281

12.4. Zarządzanie danymi w tabelach dynamicznych 286

12.4.1. Kiedy stosować tabele dynamiczne? 288

12.4.2. Kwerendowanie danych historycznych 289

Podsumowanie 291

13. Orkiestracja potoków danych 293

13.1. Orkiestracja zadań w Snowflake'u 295

13.1.1. Tworzenie schematu przechowującego
obiekty orkiestracji 296

13.1.2. Projektowanie zadań orkiestracji 297

13.1.3. Tworzenie zadań przy użyciu zależności 298

13.2. Wysyłanie powiadomień e-mail 303

13.3.	Orkiestracja za pomocą grafów zadań	304
13.3.1.	Projektowanie grafu zadań	305
13.3.2.	Tworzenie zadania głównego	306
13.3.3.	Tworzenie zadania finalizującego	307
13.3.4.	Wyświetlanie grafu zadań	308
13.4.	Monitorowanie wykonywania potoków danych	310
13.4.1.	Dodawanie funkcji dziennika zdarzeń do zadań	310
13.4.2.	Podsumowanie informacji z dzienników w powiadomieniu e-mailowym	312
13.5.	Rozwiązywanie problemów z awariami potoku danych	314
	Podsumowanie	316
14.	Kontrola integralności i kompletności danych	318
14.1.	Metody testowania danych	319
14.1.1.	Wykonywanie testów danych jako etapów w potoku	320
14.1.2.	Przeprowadzanie testów danych niezależnie od potoku	321
14.2.	Włączanie kroków testowania danych do potoku	322
14.2.1.	Tworzenie zadania sprawdzania jakości danych o kontrahentach	324
14.2.2.	Tworzenie zadania oceny jakości danych	326
14.2.3.	Uruchamianie potoku z zadaniami testowania danych	328
14.3.	Stosowanie funkcji pomiarowych Snowflake'a	329
14.3.1.	Wbudowane funkcje wskaźników danych	330
14.3.2.	Funkcje wskaźnikowe zdefiniowane przez użytkownika	331
14.3.3.	Przeglądanie szczegółów funkcji wskaźnikowych	333
14.4.	Powiadamianie użytkowników o przekroczeniu progów wskaźników	334
14.5.	Wykrywanie anomalii w woluminach danych	336
14.5.1.	Generowanie danych losowych	337
14.5.2.	Prezentowanie danych w formie wykresu liniowego w interfejsie Snowsight	339
14.5.3.	Praca z modelem wykrywania anomalii	340
	Podsumowanie	343
15.	Ciągła integracja potoków danych	345
15.1.	Rozdzielanie środowisk inżynierii danych	347
15.2.	Zarządzanie zmianami w bazie danych	348
15.2.1.	Porównanie podejścia imperatywnego i deklaratywnego w DCM	349
15.2.2.	Organizowanie kodu w repozytorium	350

15.3. Konfigurowanie Snowflake'a do współpracy z systemem Git	352
15.3.1. Tworzenie repozytorium Git	353
15.3.2. Wykonywanie poleceń z etapu repozytorium Git	356
15.4. Korzystanie z interfejsu wiersza poleceń Snowflake CLI	357
15.4.1. Instalowanie i konfigurowanie interfejsu wiersza poleceń Snowflake'a	357
15.4.2. Uruchamianie skryptów za pomocą Snowflake CLI	359
15.4.3. Ciągła integracja z wykorzystaniem interfejsu Snowflake CLI	360
15.5. Bezpieczne łączenie się z platformą Snowflake	361
15.5.1. Konfigurowanie uwierzytelniania za pomocą pary kluczy	363
15.6. Stosowanie zdobytej wiedzy w rzeczywistych scenariuszach	364
Podsumowanie	364

<i>Dodatek A. Konfiguracja środowiska Snowflake</i>	<i>367</i>
---	------------

<i>Dodatek B. Obiekty Snowflake'a wykorzystane w ćwiczeniach</i>	<i>371</i>
--	------------

Inżynieria danych w środowisku Snowflake



Zawartość rozdziału:

- wykorzystanie Snowflake'a w inżynierii danych;
- analiza obowiązków inżyniera danych w Snowflake'u;
- konstruowanie potoków danych w Snowflake'u;
- inżynieria danych w aplikacjach tworzonych z użyciem Snowflake'a.

Organizacje w niemal każdej branży gromadzą dane, często w ogromnych ilościach. Analitycy danych, danetycy i inni specjaliści biznesowi wykorzystują te informacje do uzyskiwania spostrzeżeń wspierających podejmowanie decyzji. Surowe dane rzadko nadają się do bezpośredniego wykorzystania. Użytkownicy potrzebują danych, które zostały przetworzone zgodnie z regułami biznesowymi, przekształcone tak, aby odpowiadały konkretnym potrzebom oraz ewentualnie wzbogacone o dane zewnętrzne.

Inżynieria danych to praktyka budowania rozwiązań, które wydobywają dane z systemów źródłowych, przekształcają je w użyteczne informacje i prezentują ujednolicone dane użytkownikom do dalszego wykorzystania. Akronimy **ETL** (ang. *extract, transform, load* — wyodrębnij, przekształć, załaduj) lub **ELT** (ang. *extract, load, transform* — wyodrębnij, załaduj, przekształć) są również powszechnie używane do określenia tych rozwiązań.

Inżynierowie danych są odpowiedzialni za tworzenie potoków danych, które umożliwiają analitykom, danetykom i innym użytkownikom dostęp do potrzebnych im informacji. Dostarczanie wysokiej jakości danych na czas jest kluczowe dla skutecznej analizy, dlatego inżynierowie danych odgrywają krytyczną rolę w dziedzinie analityki danych.

1.1. Snowflake a inżynieria danych

Organizacje zazwyczaj pozyskują dane z systemów źródłowych, przechowują je w hurtowni danych i udostępniają użytkownikom w postaci wymiarowych składnic danych. Dodatkowe scenariusze obejmują tworzenie jezior danych, składnic danych lub domenowych produktów danych. W każdym z tych przypadków dane są wydobywane z systemów źródłowych, wprowadzane do docelowej platformy przechowywania i przekształcane na potrzeby odbiorców końcowych.

Snowflake to nowoczesna chmurowa platforma danych służąca do przechowywania, przetwarzania i analizy danych oraz tworzenia aplikacji. Jest popularnym wyborem dla rozwiązań obejmujących hurtownie danych, jeziora danych, analizę danych czy danetykę. W poniższych punktach szczegółowo wyjaśniam możliwości środowiska Snowflake w zakresie inżynierii danych.

WSKAZÓWKA Platforma Snowflake jest nieustannie rozwijana i są do niej dodawane nowe funkcje. Najbardziej aktualne informacje są dostępne na oficjalnej stronie Snowflake'a pod adresem <https://www.snowflake.com/>.

1.1.1. Architektura Snowflake'a

Środowisko Snowflake to chmurowe rozwiązanie typu SaaS (ang. *Software as a Service* — oprogramowanie jako usługa). Obecnie wspieranymi dostawcami usług chmurowych są Amazon Web Services, Microsoft Azure i Google Cloud Platform. Gdy użytkownicy wybierają Snowflake'a jako swoją platformę danych w chmurze, nie muszą kupować, instalować, konfigurować, zarządzać ani serwisować sprzętu i oprogramowania.

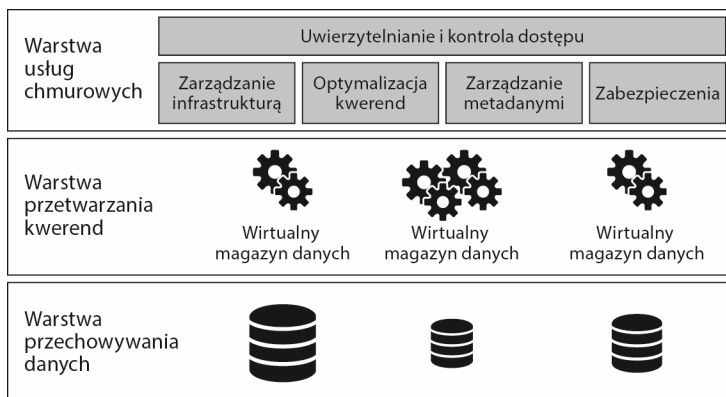
W architekturze platformy Snowflake przechowywanie i przetwarzanie danych są zarządzane oddzielnie. Dane są zapisywane jeden raz, a następnie wykorzystywane przez wiele różnych procesów, takich jak potoki inżynierii danych, silniki kwerend, rozwiązania danetyczne, udostępnianie informacji i narzędzia analityczne.

Snowflake przechowuje dane w chmurze w skompresowanym formacie kolumnowym. Zarządza on organizacją, strukturą, kompresją, metadanymi, statystykami i bezpieczeństwem danych. Użytkownicy mają dostęp do obiektów danych za pomocą kwerend SQL.

Użytkownicy wykonują kwerendy za pomocą wirtualnych magazynów danych. Każdy wirtualny magazyn wykorzystuje własne zasoby obliczeniowe, nie wpływając

na inne magazyny. Ilość zasobów dla poszczególnych magazynów można konfigurować. Snowflake może również automatycznie zwiększać lub zmniejszać przydzielone zasoby w zależności od oczekiwanej wydajności i obciążenia.

Snowflake koordynuje różne składniki za pośrednictwem warstwy usług chmurowych. Usługi te obejmują uwierzytelnianie, kontrolę dostępu, zarządzanie infrastrukturą, optymalizację kwerend, zarządzanie metadanymi oraz zabezpieczenia. Rysunek 1.1 przedstawia trzy warstwy architektury Snowflake'a: przechowywania danych, przetwarzania kwerend oraz usług chmurowych.



Rysunek 1.1. Architektura Snowflake'a składa się z trzech warstw: przechowywania danych, przetwarzania kwerend oraz usług chmurowych

Snowflake oferuje zaawansowane funkcje przechowywania danych, takie jak **klonowanie bez kopiowania** (ang. *zero-copy cloning*) i **podróże w czasie** (ang. *time travel*). Klonowanie bez kopiowania umożliwia tworzenie kopii danych bez ponoszenia dodatkowych kosztów przechowywania. Pozwala to inżynierom danych na szybkie tworzenie nowych środowisk. Na przykład mogą oni sklonować środowisko programistyczne do środowiska testowego lub skopiować środowisko produkcyjne do środowiska programistycznego w celu opracowania nowych funkcji.

Podróże w czasie mogą chronić dane przed niezamierzonymi błędami lub złośliwymi działaniami. Umożliwiają użytkownikom dostęp do danych z wybranego momentu w przeszłości i stanowią efektywny sposób tworzenia kopii zapasowych w ramach konfigurowalnego okresu.

Snowflake umożliwia replikację baz danych między kontami Snowflake'a w różnych rejonach geograficznych i platformach chmurowych. Ta funkcja zapewnia synchronizację replikowanych danych z danymi źródłowymi. Pozwala również przybliżyć dane do użytkowników w wybranym rejonie, co zmniejsza opóźnienia kwerend i służy jako kopia zapasowa.

1.1.2. Funkcje Snowflake'a używane w inżynierii danych

Snowflake oferuje różne funkcje do pozyskiwania i przekształcania danych. Funkcje te zostały opisane w poniższych podpunktach.

POBIERANIE DANYCH Z MAGAZYNÓW CHMUROWYCH

Użytkownicy mogą importować dane z plików przechowywanych w obsługiwanych dostawcach magazynów chmurowych: komorach (ang. *bucket*) AWS S3, kontenerach Azure Blob Storage oraz zasobnikach (ang. *bucket*) GCP. Pliki mogą być w formatach strukturalnych lub półstrukturalnych, takich jak CSV, JSON, XML lub Parquet. Snowflake umożliwia również pracę z plikami nieustrukturyzowanymi zawierającymi obrazy, wideo lub dźwięk. Użytkownicy mogą udostępniać adresy URL dostępu do plików oraz wczytywać te adresy URL i metadane plików do tabel Snowflake'a.

USŁUGA SNOWPIPE

Snowpipe to usługa umożliwiająca ciągle pozyskiwanie danych. Pozwala ona wprowadzać pliki z magazynu chmurowego do tabel Snowflake'a natychmiast po pojawieniu się nowych danych w chmurze. Ciągłe pozyskiwanie danych jest możliwe dzięki integracji z usługą powiadomień dostawcy chmury, która informuje Snowpipe'a o dostępności nowych plików do wczytania.

ZADANIA W SNOWFLAKE'U

Snowflake umożliwia tworzenie zadań. Zadanie realizuje instrukcje SQL, procedury składowane lub bloki skryptów Snowflake'a zgodnie z powtarzającym się harmonogramem. Zadania mogą być zależne od innych zadań, co pozwala na implementację złożonych potoków przetwarzania danych.

INTERFEJS SNOWPARK

Snowpark udostępnia biblioteki API pozwalające inżynierom danych tworzyć potoki danych w jednym z obsługiwanych języków programowania, takim jak Python, Scala czy Java. Kod napisany w tych językach jest wykonywany natywnie w środowisku Snowflake. Dzięki temu inżynierowie danych mogą budować potoki danych, korzystając z wybranego języka programowania, na przykład Pythona, i w razie potrzeby wykorzystywać dodatkowe funkcje dostępne w bibliotekach o otwartym kodzie źródłowym.

FUNKCJE ZARZĄDZANIA DANYMI

Dane przechowywane w środowisku Snowflake mogą zawierać wrażliwe informacje, które powinny być dostępne tylko dla uprawnionych użytkowników. W architekturach wielotenantowych dane różnych organizacji lub jednostek biznesowych są przechowywane w jednej bazie danych. Mimo to dostęp do danych

jest ograniczony do użytkowników z tej samej organizacji. Aby obsłużyć takie przypadki użycia, Snowflake oferuje funkcje zarządzania danymi, takie jak reguły maskowania i reguły dostępu do wierszy, które uniemożliwiają nieupoważnionym użytkownikom dostęp do wrażliwych lub poufnych danych.

SNOWFLAKE MARKETPLACE

Snowflake Marketplace to platforma, na której użytkownicy mogą odkrywać i uzyskiwać dostęp do zestawów danych udostępnionych przez inne organizacje. Mogą wybierać interesujące ich zestawy danych i wykorzystywać je do wzbogacania własnych analiz. Mają również możliwość publikowania własnych zestawów danych, aby udostępnić je innym. Zestawy danych można udostępniać publicznie lub prywatnie, co umożliwia bezpieczne dzielenie się danymi i współpracę zarówno wewnątrz organizacji, jak i z podmiotami zewnętrznymi.

SNOWFLAKE PARTNER CONNECT

Dzięki bogatemu środowisku partnerów biznesowych Snowflake płynnie integruje się z wieloma popularnymi narzędziami i usługami zewnętrznymi z różnych kategorii, takich jak analityka biznesowa, integracja danych, uczenie maszynowe, danetyka, bezpieczeństwo, zarządzanie danymi i obserwowalność. Użytkownicy wykorzystują te narzędzia do pozyskiwania danych w procesach inżynierii danych, przeprowadzania analiz, tworzenia raportów i pulpitu oraz pracy z algorytmami uczenia maszynowego i danetycznymi

Dzięki Snowflake'owi inżynierowie danych mogą korzystać z wielu funkcji do budowania potoków danych, dostosowując je do swoich preferencji, założeń architektonicznych i potrzeb biznesowych.

1.2. Obowiązki inżyniera danych w Snowflake'u

Inżynierowie danych są odpowiedzialni za projektowanie, budowanie, orkiestrację, optymalizację, monitorowanie i utrzymywanie potoków, które pozyskują surowe dane z systemów źródłowych i przekształcają je w użyteczne informacje gotowe do dalszego wykorzystania. Dbają oni o to, by dane były łatwo dostępne, bezpieczne, wiarygodne i osiągalne dla użytkowników zgodnie z ich potrzebami.

Rolę inżyniera danych zwykle pełnią specjaliści z wcześniejszym doświadczeniem, ponieważ wymaga ona szerokiego zakresu umiejętności. Wielu inżynierów danych ma za sobą pracę w obszarze tworzenia oprogramowania lub analizy biznesowej i stopniowo rozszerza swoje kompetencje. Pracując w Snowflake'u, zdobywają oni wiedzę i doświadczenie związane z funkcjami tej platformy.

Znajomość chmury obliczeniowej jest kluczowa podczas korzystania ze Snowflake'a. Na przykład inżynierowie danych często pracują z chmurowymi magazynami obiektów i wiedzą, jak bezpiecznie łączyć się ze środowiskami chmurowymi. Wiele organizacji przechodzi na rozwiązania chmurowe i przenosi swoją

infrastrukturę do dostawców usług w chmurze, co oznacza, że inżynierowie danych coraz swobodniej poruszają się w środowisku chmurowym.

Inżynierowie danych znają organizację przechowywania danych w środowisku Snowflake, co pozwala im pisać wydajne kwerendy SQL pobierające i przetwarzające dane. Rozumieją, w jaki sposób dane są przechowywane w mikropartycjach Snowflake'a oraz kiedy i jak stosować klasteryzację, optymalizację wyszukiwania, przyspieszanie kwerend lub inne funkcje w celu poprawy wydajności kwerend.

Znają oni również sposoby efektywnego kosztowo wykorzystania wirtualnych magazynów Snowflake'a i rozumieją, jak ich rozwiązania wpływają na zużycie mocy obliczeniowej. Wirtualne magazyny mogą także poprawić wydajność kwerend, ponieważ utrzymują pamięć podręczną danych, którą inne kwerendy wykonywane w tym samym magazynie mogą ponownie wykorzystać.

Podczas tworzenia rozwiązań w środowiskach korporacyjnych inżynierowie danych stosują, tam gdzie jest to możliwe, zasady inżynierii oprogramowania, takie jak najlepsze praktyki pisanie kodu, zautomatyzowane testy, kontrola wersji kodu, ciągła integracja oraz ciągłe wdrażanie. Jeśli korzystają z interfejsu Snowpark, posiadają również znajomość języka programowania takiego jak Python.

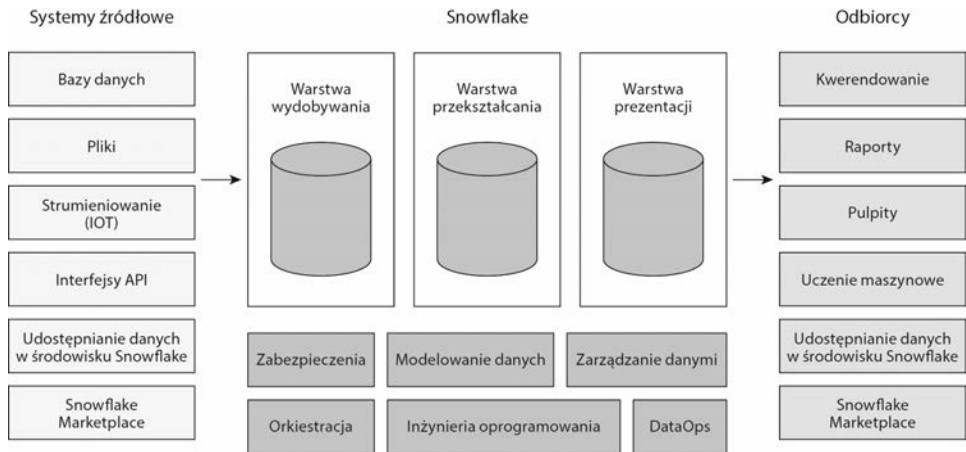
Umiejętność bezpiecznego łączenia się z platformą Snowflake ma kluczowe znaczenie dla uniknięcia naruszeń bezpieczeństwa i przypadkowego udostępnienia poufnych danych osobom nieupoważnionym. Każdy inżynier danych musi koniecznie i szczegółowo znać różne metody nawiązywania połączenia ze środowiskiem Snowflake.

Rysunek 1.2 przedstawia ogólne elementy składowe inżynierii danych służące do budowy potoków danych w ramach platformy chmurowej Snowflake. Nie wszystkie składniki są istotne w każdym projekcie. W zależności od rodzaju tworzonego rozwiązania inżynierowie danych wykorzystują tylko te elementy, które są odpowiednie i niezbędne w danym środowisku. Głównymi komponentami inżynierii danych są:

- pozyskiwanie danych z systemów źródłowych i wprowadzanie ich do Snowflake'a,
- wykonywanie przekształceń danych istotnych dla rozwiązania;
- prezentowanie danych odbiorcom końcowym do celów analityki, tworzenia raportów, danetyki lub innych zastosowań.

Oprócz głównych składników inżynierowie danych wykorzystują również komponenty bazowe podczas budowania potoków przetwarzania danych. Te bazowe elementy obejmują:

- zabezpieczenia,
- modelowanie danych,
- zarządzanie danymi,



Rysunek 1.2. Elementy składowe inżynierii danych w Snowflake'u obejmują wydobywanie danych z systemów źródłowych, wprowadzanie ich na Snowflake'a, wykonywanie przekształceń oraz udostępnianie danych odbiorcom końcowym. Proces ten uwzględnia podstawowe elementy i najlepsze praktyki związane z bezpieczeństwem, modelowaniem danych, zarządzaniem danymi, orkiestracją, inżynierią oprogramowania oraz metodologią DataOps

- inżynierię oprogramowania,
- orkiestrację,
- metodologię DataOps.

Każdy z tych elementów został szczegółowo opisany w kolejnych podrozdziałach (zobacz rysunek 1.2).

1.2.1. Pozyskiwanie danych z systemów źródłowych

Pierwszym krokiem w procesie inżynierii danych jest pozyskiwanie danych z systemów źródłowych. Inżynierowie danych zdobywają podstawową wiedzę na temat biznesowego znaczenia danych i ich wykorzystania w rozwiązaniach końcowych. Często kilka systemów źródłowych dostarcza dane do rozwiązania, a inżynierowie danych tworzą potoki, które płynnie pobierają te dane na jedną platformę, tworząc ujednolicony model dla użytkowników.

Inżynierowie danych zapoznają się z technologiami używanymi w systemach źródłowych — na przykład sprawdzają, czy zestawy danych istnieją jako tabele w bazach danych, czy pliki w chmurze, czy są generowane jako wynik zapytań API z aplikacji, czy pochodzą ze strumieni danych z urządzeń IoT, czy też z innych źródeł. Identyfikują najlepsze sposoby dostępu do danych i przygotowania ich do umieszczenia w Snowflake'u. Dane mogą być dostępne nie tylko z zewnętrznych źródeł, lecz również poprzez bezpieczne udostępnianie danych Snowflake'a lub Snowflake Marketplace.

PRZYROSTOWE POBIERANIE DANYCH

Potoki danych zazwyczaj pobierają dane z systemów źródłowych przyrostowo, przetwarzając tylko nowe i zmienione informacje. Inżynierowie danych współpracują z właścicielami systemów źródłowych, aby ustalić najlepszą metodę wykrywania zmian. Wykorzystują filtrowanie według znaczników czasu, mechanizmy śledzenia zmian w bazach danych lub inne techniki. W przypadku pobierania danych z przestarzałych systemów, których właściciele są niedostępni, inżynierowie danych czasami muszą sami określić najlepsze podejście do wykrywania zmian, opierając się na własnym doświadczeniu i wiedzy specjalistycznej.

Podczas pracy z systemami, które nie umożliwiają łatwej identyfikacji zmian w danych, inżynierowie mogą wykorzystać *strumienie* Snowflake'a. Strumień śledzi operacje takie jak wstawianie, aktualizowanie czy usuwanie, wykonywane na danych źródłowych między dwoma punktami w czasie.

NARZĘDZIA A KORZYSTANIE Z WŁASNEGO KODU

Rozpoczynając nowy projekt z zakresu inżynierii danych, inżynierowie danych współpracują z architektami rozwiązań i innymi zainteresowanymi stronami, aby wybrać najlepsze narzędzia i technologie do wydobywania danych z systemów źródłowych do środowiska Snowflake. Mogą to być narzędzia o otwartym kodzie źródłowym lub komercyjne, a także własny kod napisany z wykorzystaniem bibliotek i technologii otwartych lub własnościowych.

Podczas tworzenia własnego kodu inżynierowie danych rozumieją formaty danych źródłowych oraz znają odpowiednie polecenia i opcje do wprowadzania tych danych do Snowflake'a. Są również zaznajomieni z podstawowymi komentarami i opcjami Snowflake'a używanymi przez zewnętrzne narzędzia do ekstrakcji danych. Wszystkie narzędzia korzystają z tych poleceń w tle, a inżynierowie danych powinni być w stanie rozwiązywać ewentualne problemy.

1.2.2. Przekształcanie danych

Dane pobrane z systemów źródłowych są przekształcane z ich surowego formatu w użyteczne i przydatne informacje zgodnie z wymaganiami użytkowników. Inżynierowie danych zazwyczaj otrzymują reguły transformacji od analityków danych, architektów systemów lub czasami bezpośrednio od użytkowników. W mniejszych zespołach inżynierii danych inżynierowie mogą współpracować z użytkownikami biznesowymi w celu zdefiniowania zasad przekształcania danych.

NARZĘDZIA DO PRZEKSZTAŁCANIA DANYCH

Narzędzia używane przez inżynierów danych do przekształcania danych są rozmaite. Mogą one przeprowadzać transformacje za pomocą kwerend SQL, procedur składowanych lub funkcji. Niektóre narzędzia do pobierania danych z systemów źródłowych również obsługują przekształcenia danych i mogą być wykorzystywane w tym celu.

Inżynierowie danych mogą też pisać kod w usłudze Snowpark, co umożliwia im tworzenie przekształceń w jednym z obsługiwanych języków programowania, takim jak Python, Java czy Scala. Mogą także korzystać z popularnych otwartoźródłowych bibliotek do przekształcania danych.

SPRAWDZANIE POPRAWNOŚCI DANYCH

Oprócz transformacji danych inżynierowie danych dodają do potoków etapy walidacji, dzięki czemu odbiorcy mogą mieć pewność, że otrzymywane dane są wiarygodne i spójne. Potoki danych często zawierają techniczne reguły walidacji, takie jak sprawdzanie naruszeń kluczy głównych lub obcych, ponieważ Snowflake nie wymusza ograniczeń kluczy głównych ani obcych w standardowych tabelach. Odbiorcy mogą definiować dodatkowe reguły walidacji danych. Dane, które nie spełniają tych reguł, mogą zostać oznaczone lub wykluczone z dalszych analiz.

1.2.3. Udostępnianie danych odbiorcom końcowym

Po pobraniu danych z systemów źródłowych i odpowiednim ich przekształceniu są one prezentowane odbiorcom w postaci dostosowanej do narzędzi i rozwiązań używanych w dalszych etapach przetwarzania. Inżynierowie danych zaspokajają potrzeby biznesowe odbiorców, dostarczając dane zgodnie z oczekiwaniami dotyczącymi terminowości i wydajności.

Inżynierowie danych tworzą odpowiednie interfejsy użytkownika tam, gdzie są one potrzebne, takie jak API, jeśli ma to zastosowanie, lub prezentują dane w modelu odpowiednim dla narzędzia używanego do analizy. Dostarczają również metadane, aby ułatwić odnajdywanie danych, na przykład przez udostępnianie komentarzy do tabel i atrybutów w bazie danych oraz innych istotnych informacji, takich jak między innymi czas ostatniej aktualizacji danych, liczba zaimportowanych rekordów i wskaźniki jakości danych (jeśli są stosowane).

W przeciwieństwie do wielu tradycyjnych baz danych kwerendy SQL w środowisku Snowflake nie wymagają ręcznej optymalizacji, która zwykle pochłania sporo czasu administratorów. Tutaj większość kwerend działa wydajnie bez dodatkowych zabiegów. Niemniej zdarzają się sytuacje, w których administratorzy i inżynierowie danych mogą poprawić wydajność kwerend o wysokiej złożoności, operujących na dużych zestawach danych lub zużywających zbyt wiele zasobów. W takich przypadkach inżynierowie mogą zwiększyć wydajność, wykorzystując funkcje Snowflake'a takie jak klasteryzacja, optymalizacja wyszukiwania, widoki zmaterializowane czy tabele dynamiczne.

1.2.4. Stosowanie podstawowych składników

Choć inżynieria danych zasadniczo obejmuje pozyskiwanie danych z systemów źródłowych, wprowadzanie ich do środowiska Snowflake i przekształcanie na potrzeby dalszego wykorzystania, to inżynierowie danych stosują również dodatkowe komponenty w zależności od potrzeb.

ZABEZPIECZENIA

Niezależnie od tego, czy rozwiązania IT są przechowywane w chmurze, czy lokalnie, mogą być narażone na zagrożenia i luki w zabezpieczeniach. Ponieważ Snowflake działa w chmurze, konieczne są dodatkowe środki ostrożności, aby zapobiec nieautoryzowanemu dostępowi. Administratorzy Snowflake'a zazwyczaj zabezpieczają dostęp na różne sposoby, takie jak ograniczenie adresów IP, z których użytkownicy mogą łączyć się z platformą, lub wymaganie silnych metod uwierzytelniania, na przykład za pomocą uwierzytelniania wieloskładnikowego.

Podczas tworzenia potoków danych inżynierowie nieuchronnie w pewnym momencie łączą się z platformą Snowflake. Wszelkie dane uwierzytelniające używane do połączenia z nią są bezpiecznie przechowywane w skrytkach dostawcy chmury lub lokalnie w bezpiecznym miejscu, które nigdy nie jest udostępniane innym użytkownikom ani zapisywane w repozytorium, gdzie inni mogliby je zobaczyć. Inżynierowie danych muszą działać odpowiedzialnie, aby zagwarantować, że dane dostępne do konta Snowflake'a nie zostaną celowo ani przypadkowo udostępnione komukolwiek innemu.

Użytkownicy powinni mieć minimalne uprawnienia niezbędne do wykonania zadania. Nawet gdy inżynierom danych przyznaje się bardziej zaawansowane role administracyjne w Snowflake'u do wykonywania określonych zadań, nie powinni oni korzystać z tych ról domyślnie. Powinni używać swoich wyznaczonych ról programistycznych i świadomie przełączać się na bardziej zaawansowane role administracyjne tylko wtedy, gdy jest to konieczne do wykonania konkretnego zadania administracyjnego.

Inżynierowie danych rozumieją zasady kontroli dostępu opartej na rolach w Snowflake'u i używają odpowiednich ról podczas tworzenia obiektów w tym systemie. Choć niekoniecznie muszą być odpowiedzialni za konfigurację mechanizmu kontroli dostępu, to powinni znać jego strukturę, aby móc prawidłowo z niego korzystać.

ZARZĄDZANIE DANYMI

Większość organizacji posiada zasady zarządzania danymi, które ograniczają nieupoważnionym użytkownikom dostęp do poufnych lub wrażliwych informacji. Czasami, szczególnie w środowiskach o wysokim poziomie bezpieczeństwa i ścisłej regulacji, inżynierowie danych nie otrzymują dostępu do rzeczywistych danych produkcyjnych, lecz pracują na próbkach danych, na danych z zamaskowanymi informacjami wrażliwymi lub na podzbiorach, z których usunięto poufne informacje.

Inżynierowie danych współpracują z architektami danych i właścicielami domen, aby zrozumieć wymagania dotyczące zarządzania danymi. Przestrzegają tych wymagań podczas budowania potoków danych. Na przykład mogą przechowywać wrażliwe atrybuty w oddzielnych tabelach z bardziej restrykcyjnymi regułami dostępu.

MODELOWANIE DANYCH

W zależności od wielkości organizacji, w której pracują inżynierowie danych, może przypaść im obowiązek modelowania danych. W dużych organizacjach często istnieją wyspecjalizowane zespoły lub eksperci odpowiedzialni za tworzenie modeli danych. W mniejszych inżynierowie danych mogą modelować dane w zakresie swoich obowiązków.

Niezależnie od tego, kto tworzy model danych, inżynier danych musi rozumieć koncepcje modelowania, aby poprawnie wczytywać dane do docelowych modeli. Na przykład modele mogą wymagać, aby potoki danych generowały klucze zastępcze podczas procesu wczytywania danych.

Oto przykłady standardowych modeli danych stosowanych w rozwiązaniach hurtowni danych i analityce:

- modele relacyjne dla korporacyjnych hurtowni danych;
- modele zespołów danych, takie jak skarbcze danych (ang. *data vault*) czy modelowanie kotwicowe (ang. *anchor modeling*), które są szczególnie przydatne jako korporacyjne hurtownie danych przy pobieraniu informacji z wielu źródeł;
- modele danych wymiarowych dla składnic danych i narzędzi do tworzenia raportów.

Niektóre modele danych umożliwiają stosowanie powtarzalnych wzorców ładowania, które automatyzują proces wprowadzania danych. Inżynierowie danych implementują takie wzorce, co pomaga im tworzyć bardziej ustandaryzowane i wydajne potoki danych.

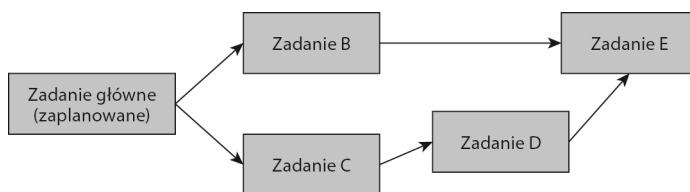
ORKIESTRACJA

Większość potoków danych działa według harmonogramu ustalonego z właściwościami systemów źródłowych i odbiorcami końcowymi. W zależności od potrzeb użytkowników potoki mogą być uruchamiane masowo albo w mikropartiach, na przykład codziennie, co godzinę lub co kilka minut.

Gdy system zawiera wiele potoków danych, inżynierowie danych projektują zależności między potokami i ich składnikami. Mogą oni wykorzystywać zewnętrzne narzędzia do koordynowania przepływu danych, polecenia systemu operacyjnego lub zadania Snowflake'a do planowania potoków w uporządkowany sposób.

Zadania w Snowflake'u umożliwiają tworzenie skierowanego grafu acyklicznego (ang. *Directed Acyclic Graph* — DAG), który zawiera serię powiązanych ze sobą zadań. Każdy DAG zaczyna się od jednego zadania głównego, a zadania zależne są wykonywane w tym samym kierunku bez tworzenia pętli. Rysunek 1.3 przedstawia przykładowy DAG zadań w Snowflake'u.

Inżynierowie danych monitorują zaplanowane potoki, aby upewnić się, że działają one bez błędów i poprawnie ładują dane. Podczas projektowania potoków danych uwzględniają mechanizmy rejestrowania zdarzeń, które umożliwiają



Rysunek 1.3. Skierowany graf acykliczny (DAG) zadań. Graf składa się z zadania głównego oraz zadań zależnych, które rozpoczynają się po zakończeniu poprzednich zadań

skuteczne rozwiązywanie problemów w przypadku wystąpienia błędów. Projektują również potoki danych w taki sposób, aby możliwe było ponowne uruchomienie całego procesu od początku lub wznowienie go od miejsca, w którym wystąpił błąd.

INŻYNIERIA OPROGRAMOWANIA

Inżynieria oprogramowania stanowi znaczną część pracy inżynierów danych. Tworzą oni, testują i utrzymują potoki przetwarzania danych, które często są rozbudowanymi systemami składającymi się z kodu programistycznego. Nawet jeśli większość tego kodu to kwerendy SQL, a nie tradycyjne konstrukcje języków programowania, nadal stosuje się zasady inżynierii oprogramowania, takie jak konwencje nazewnictwa, modularyzacja kodu, czytelność kodu, wersjonowanie i łatwość utrzymania systemu.

Inżynierowie danych tworzą również wydajne, solidne i niezawodne systemy. Weźmy na przykład potok danych, który zostałby przypadkowo uruchomiony dwukrotnie. W takiej sytuacji nie powinno to prowadzić do duplikowania danych, a potok nie powinien zostać uruchomiony, gdy poprzednia instancja tego samego potoku jest wciąż w trakcie przetwarzania.

Istotnym elementem inżynierii danych jest zrozumienie i wykorzystanie różnych środowisk w cyklu życia oprogramowania. Organizacje zazwyczaj korzystają ze środowiska programistycznego, w którym programiści pracują nad swoim kodem. Następnie kod jest wdrażany do środowiska testowego, gdzie sprawdza się, czy utworzone potoki danych działają zgodnie z oczekiwaniami. Niektóre duże firmy mogą również posiadać środowiska do testów akceptacyjnych. Po przetestowaniu i zatwierdzeniu kodu jest on wdrażany do środowiska produkcyjnego.

DATAOPS

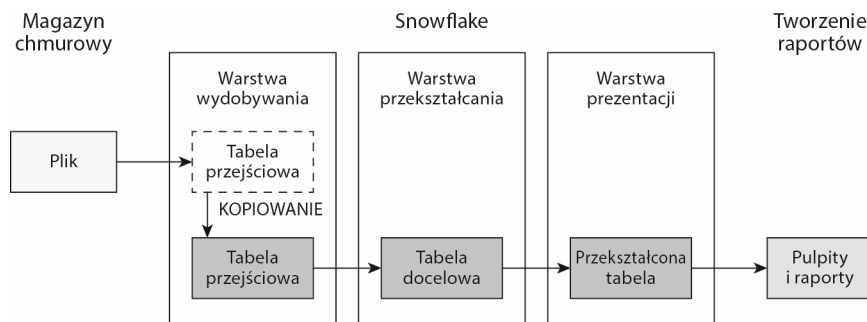
DataOps to ogólne pojęcie wywodzące się z metodyki zwinnego tworzenia oprogramowania, które ma zastosowanie w całym cyklu życia danych, obejmującym wydobywanie, pozyskiwanie, przekształcanie i prezentowanie. Łączy ona procesy i technologie, które poprawiają jakość oraz umożliwiają ciągły rozwój i integrację potoków inżynierii danych.

DataOps, podobnie jak DevOps w inżynierii oprogramowania, zapewnia ciągłe dostarczanie działającego oprogramowania poprzez wykorzystanie narzędzi i technologii umożliwiających pracę zespołową. Na przykład zespoły inżynierii danych korzystają z repozytoriów kodu, takich jak Git, zintegrowanych z procesami CI/CD. Mogą one również automatyzować testy kodu potoków danych i walidację samych danych. Nawet w małych organizacjach, gdzie pracuje tylko jeden inżynier danych, praktyki DataOps pomagają lepiej zorganizować i ustrukturyzować pracę, zapewniając spójne i ciągłe wdrażanie do środowiska produkcyjnego.

1.3. Budowanie potoków danych

Potok danych (ang. *data pipeline*) to zestaw składników przetwarzających dane, który przenosi je ze źródła do miejsca docelowego. W zależności od wymagań potok może również wykonywać przekształcanie danych. Każdy potok danych jest inny. Na przykład inżynier danych pobiera pliki z chmurowego magazynu obiektów do środowiska Snowflake i udostępnia dane z tych plików odbiorcom końcowym, którzy tworzą na ich podstawie raporty.

Rysunek 1.4 przedstawia przykładowy potok danych. Pliki z chmurowego magazynu obiektów są udostępniane Snowflake'owi poprzez zewnętrzne miejsce przechowywania i wczytywane do **tabeli przejściowej** (ang. *staging table*). Następnie potok ładuje dane do tabeli docelowej, przeprowadza niezbędne transformacje i umieszcza przetworzone dane w odpowiedniej tabeli. Na koniec udostępnia te dane do dalszego raportowania, kwerendowania lub analizowania przez użytkowników końcowych.



Rysunek 1.4. Potok danych, który pobiera plik danych z chmurowego dostawcy pamięci masowej, wczytuje dane do tabeli Snowflake'a i przekształca je na potrzeby tworzenia raportów

Przykład przedstawiony na rysunku 1.4 to jeden z najczęściej spotykanych potoków inżynierii danych w środowisku Snowflake. Wykorzystuje on miejsce zewnętrzne wskazujące na lokalizację w chmurowym magazynie obiektów. Następnie używa polecenia COPY do wczytania danych z magazynu chmurowego do

tabeli przejściowej Snowflake'a. Potem dane te są ładowane do tabeli docelowej poprzez dołączanie ich za pomocą polecenia `INSERT` lub scalanie za pomocą polecenia `MERGE`. Tabela docelowa jest zwykle projektowana, jeszcze zanim inżynier danych zbuduje potok danych jako część modelu danych warstwy przekształcania.

Wszelkie niezbędne transformacje są przeprowadzane za pomocą instrukcji `SQL` lub procedur składowanych, a przetworzone dane są zapisywane w przeznaczonych do tego tabeli. Na koniec dane są prezentowane w formacie odpowiednim do sporządzania raportów. Zamiast budować etapy transformacji danych w potokach, inżynierowie danych mogą tworzyć widoki `SQL` do przekształcania danych, ponieważ widoki są łatwe w utrzymaniu i nie wymagają ponownego ładowania danych. Jednak widoki mogą nie zapewniać wydajności wymaganej przez użytkowników. W takich sytuacjach inżynier danych decyduje o najlepszym podejściu do poprawy wydajności, często poprzez materializację danych.

W praktyce stosuje się wiele rodzajów potoków danych wykorzystujących różne składniki. Omówię je dokładniej w kolejnych rozdziałach.

1.4. Inżynieria danych w zastosowaniach Snowflake'a

Najczęstsze zastosowania inżynierii danych z wykorzystaniem platformy Snowflake obejmują jeziora danych, hurtownie danych, składnice danych, siatki danych, analizę biznesową, danetykę oraz dogenerowywanie danych przy użyciu dużych modeli językowych.

JEZIORO DANYCH

Jezioro danych (ang. *data lake*) to repozytorium przechowujące duże ilości danych o pełnej, częściowej lub zerowej strukturze. Snowflake umożliwia przechowywanie danych w różnych typach tabel, w tym w tabelach standardowych, zewnętrznych, hybrydowych i innych. Wiele równoczesnych zadań może uzyskiwać dostęp do danych bez rywalizacji o zasoby. Użytkownicy mogą odpytywać dane relacyjne lub częściowo ustrukturyzowane w jeziorze danych za pomocą składni `SQL` lub innych języków obsługiwanych w środowisku programistycznym Snowpark.

HURTOWNIA DANYCH

Hurtownia danych (ang. *data warehouse*) to korporacyjne repozytorium danych gromadzonych z jednego bądź wielu systemów źródłowych lub zewnętrznych źródeł, takich jak Snowflake Marketplace. Głównym celem hurtowni danych jest dostarczanie ujednoliconych danych historycznych do tworzenia raportów, pulpitów i analiz dla użytkowników biznesowych. Umożliwia to użytkownikom

biznesowym wyciąganie z tych danych wniosków, które stanowią podstawę do podejmowania decyzji biznesowych.

SKŁADNICA DANYCH

Składnice danych (ang. *data mart*) przechowują informacje dla pojedynczych działów biznesowych, takich jak marketing czy finanse. Działy te mogą szybko i łatwo uzyskać dostęp do danych w składnicach, zamiast przeszukiwać całą hurtownię danych. Ponieważ budowa pełnej hurtowni danych może pochłonąć znaczne zasoby i zająć dużo czasu, niektóre organizacje decydują się na tworzenie wyłącznie składnic bez budowania centralnej hurtowni danych.

SIATKA DANYCH

Siatka danych (ang. *data mesh*) to stosunkowo nowe podejście do analizy danych, w którym zamiast monolitycznej hurtowni danych poszczególne zespoły dziedzinowe tworzą własne produkty danych. Często przypominają one składnice danych. Dzięki lepszym narzędziom oraz łatwiejszym w obsłudze platformom i technologiom użytkownicy biznesowi mogą samodzielnie przygotowywać i eksplorować dane. Na przykład mogą oceniać jakość danych w swoich jednostkach biznesowych. Mogą również wykonywać proste przekształcenia. W rezultacie część zadań związanych z inżynierią danych przechodzi od inżynierów danych na użytkowników biznesowych.

Nawet jeśli zadania związane z inżynierią danych zostaną przekazane użytkownikom biznesowym, zapotrzebowanie na wykwalifikowanych inżynierów danych nie zniknie. Wiele technicznych aspektów potoków danych wymaga solidnej wiedzy z zakresu inżynierii oprogramowania, której użytkownicy biznesowi mogą nie posiadać ze względu na inny obszar specjalizacji. Inżynierowie danych nadal odgrywają istotną rolę w podejściu opartym na siatce danych, jednak ich rola może się przenieść z działu IT do odpowiednich działów biznesowych.

Snowflake może przechowywać produkty danych, a funkcja udostępniania danych w niej umożliwia ekspozycję produktów danych w całej organizacji.

ANALITYKA BIZNESOWA

Analityka biznesowa (ang. *business intelligence*) to proces analizowania danych i uzyskiwania na ich podstawie informacji umożliwiających podejmowanie decyzji. Rozwiązania z zakresu analityki biznesowej zazwyczaj czerpią dane z firmowych hurtowni danych, składnic danych, produktów danych lub kombinacji tych źródeł.

DANETYKA

Danetyka (ang. *data science*) to szeroka dziedzina obejmująca zastosowania takie jak zaawansowana analityka, sztuczna inteligencja i uczenie maszynowe. Zasadniczo danetyka ma na celu odkrywanie użytecznych informacji ze źródeł

danych, które są wykorzystywane do podejmowania decyzji i planowania strategicznego w organizacjach. Usługa Snowpark umożliwia danetykom budowanie modeli przy użyciu wybranego języka programowania, szczególnie Pythona, który ma bogaty wybór bibliotek do uczenia maszynowego. Na przykład danetycy mogą tworzyć, trenować i wdrażać modele uczenia maszynowego w środowisku Snowpark, korzystając z popularnych bibliotek takich jak scikit-learn, TensorFlow czy PyTorch. Wraz z wprowadzeniem środowiska Snowflake Cortex niektóre funkcje uczenia maszynowego, takie jak prognozowanie szeregów czasowych, wykrywanie anomalii czy klasyfikacja, są dostępne bezpośrednio z poziomu poleceń SQL w usłudze Snowflake.

DOGENEROWYWANIE DANYCH PRZY UŻYCIU DUŻYCH MODELI JĘZYKOWYCH

Generatywna sztuczna inteligencja to technologia, która tworzy tekst, obrazy, kod programistyczny i inne rodzaje treści. **Duże modele językowe** (ang. *Large Language Models* — LLM) to podkategoria generatywnej sztucznej inteligencji specjalizująca się w tworzeniu tekstu. Wiele zastosowań biznesowych, takich jak klasyfikacja tekstu, analiza dokumentów, ocena nastroju czy tłumaczenie maszynowe, opiera się na LLM-ach. Wyniki generowane przez LLM-y mogą wzbogacać dane w różnych rozwiązaniach platformy Snowflake.

Programiści mogą korzystać z funkcji zewnętrznych modeli językowych w usłudze Snowpark, wywołując odpowiednie API i zapisując wyniki w tabelach Snowflake'a do dalszej analizy. Wraz z wprowadzeniem środowiska Snowflake Cortex wybrane funkcje LLM-ów, takie jak ocena nastroju, uzupełnianie podpowiedzi, streszczanie, tłumaczenie, wydobywanie odpowiedzi z nieustrukturyzowanego tekstu czy osadzanie wektorowe stały się dostępne bezpośrednio z poziomu poleceń SQL w usłudze Snowflake.

Podsumowanie

- Snowflake to nowoczesna platforma danych w chmurze, która doskonale nadaje się do rozwiązań wymagających intensywnego przetwarzania danych, takich jak jeziora danych, hurtownie danych, składnice danych, siatki danych, analityka biznesowa, danetyka oraz dogenerowywanie danych przy użyciu modeli językowych.
- Inżynierowie danych w środowisku Snowflake są odpowiedzialni za projektowanie, tworzenie, koordynowanie, optymalizowanie, monitorowanie i utrzymywanie potoków danych. Potoki te przetwarzają surowe dane z systemów źródłowych, przekształcając je w użyteczne informacje gotowe do dalszego wykorzystania.

- Pierwszym krokiem w procesie przetwarzania danych jest ich pozyskiwanie z systemów źródłowych. Inżynierowie danych uzyskują dostęp do tych systemów i wprowadzają dane do Snowflake'a, korzystając z własnoręcznie napisanych poleceń lub używając specjalistycznych narzędzi.
- Dane pobrane z systemów źródłowych na Snowflake'a są przekształcane ze stanu surowego do postaci i formatu odpowiednich dla odbiorców końcowych.
- Transformacje danych można przeprowadzać za pomocą narzędzi, instrukcji SQL lub procedur składowanych, a także kodować w usłudze Snowpark przy użyciu jednego z obsługiwanych języków programowania.
- Dane dla odbiorców końcowych są prezentowane w formie, która ułatwia dostęp do nich i efektywne ich wykorzystanie w narzędziach raportowych lub innych rozwiązaniach, takich jak modele uczenia maszynowego w środowisku Snowpark.
- Potok danych to zestaw etapów przetwarzania, który przenosi dane ze źródła do miejsca docelowego, wykonując niezbędne transformacje po drodze.
- Podczas tworzenia potoków danych dla Snowflake'a inżynierowie danych wykorzystują podstawowe komponenty, takie jak zabezpieczenia, zarządzanie danymi, modelowanie danych, orkiestracja, inżynieria oprogramowania oraz metodyka DataOps.
- Najczęstsze zastosowania inżynierii danych z wykorzystaniem Snowflake'a to jeziora danych, hurtownie danych, składnice danych, siatki danych, analityka biznesowa, danetyka oraz dogenerowywanie danych przy użyciu modeli językowych.

PROGRAM PARTNERSKI

— GRUPY HELION —



1. ZAREJESTRUJ SIĘ
2. PREZENTUJ KSIĄŻKI
3. ZBIERAJ PROWIZJĘ

Zmień swoją stronę WWW w działający bankomat!

Dowiedz się więcej i dołącz już dzisiaj!

<http://program-partnerski.helion.pl>

GRUPA
Helion

Snowflake

➤ **Pozycja obowiązkowa, jeżeli chcesz osiągnąć wyżyny w dziedzinie inżynierii danych!**

— **Isabella Renzetti**, ekspertka z dziedziny analizy danych

Snowflake to kompleksowa platforma chmurowa do przechowywania i analizy danych, oferująca niemal nieograniczoną skalowalność i szybkie, elastyczne usługi obliczeniowe. Umożliwia tworzenie i rozwijanie potoków danych, a jej nowe funkcjonalności, takie jak wyszukiwanie wektorowe, automatyczne konwersje tekstu do SQL czy generowanie kodu, korzystają z technik AI. Jeśli pracujesz z danymi, Snowflake otwiera przed Tobą zupełnie nowe możliwości.

➤ **Niesamowita książka! Przejdiesz od zera do środowiska produkcyjnego Snowflake!**

— **Doyle Turner**, Microsoft

Z tą książką krok po kroku będziesz rozwijać umiejętności potrzebne do wykonywania codziennych zadań z zakresu inżynierii danych w Snowflake. Utworzysz swój pierwszy potok, a potem będziesz go rozbudowywać o coraz bardziej zaawansowane funkcje: zarządzanie danymi i bezpieczeństwem, integrację CI/CD czy wzbogacanie danych przy użyciu generatywnej AI. To praktyczny poradnik pełen kodu, przykładów i wskazówek, które pozwolą Ci w krótkim czasie wejść na zupełnie nowy poziom pracy z danymi.

➤ **Wyczerpująca, aktualna i wypełniona praktycznymi fragmentami kodu!**

— **Albert Nogués**, Danone

W książce:

- wprowadzanie danych z usług chmurowych za pomocą API lub ze Snowflake Marketplace
- orkiestracja potokami danych za pomocą strumieni i zadań
- optymalizacja wydajności i kosztów
- projektowanie mechanizmów kontroli dostępu i ochrony danych
- wdrażanie ciągłe oraz ciągła integracja obiektów i kodu Snowflake
- wzbogacanie danych za pomocą generatywnej AI

➤ **Mistrzostwo! Odkryjesz rzeczywisty potencjał platformy Snowflake!**

— **Shankar Narayanan**, Microsoft

Helion 		KOD KORZYŚCI Sięgnij po więcej! 	
 helion.pl			
 HELION S.A. ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl	ISBN 978-83-289-2898-5  9 788328 928985		
Cena: 99,00 zł			

Maja Ferle jest architektką danych z ponad 30-letnim stażem w dziedzinie analizy danych, hurtowni danych, analityki biznesowej, inżynierii danych, modelowania danych i administracji bazami danych. Posiada certyfikaty SnowPro Advanced Data Engineer i SnowPro Advanced Data Analyst. Jest również ekspertką Snowflake (uzyskała certyfikaty SnowPro Subject Matter Expert i Snowflake Data Superhero).