



IJIRCCE

e-ISSN: 2320-9801 | p-ISSN: 2320-9798



INTERNATIONAL JOURNAL OF INNOVATIVE RESEARCH

IN COMPUTER & COMMUNICATION ENGINEERING

Volume 12, Issue 1, January 2024

ISSN INTERNATIONAL
STANDARD
SERIAL
NUMBER
INDIA

Impact Factor: 8.379



9940 572 462



6381 907 438



ijircce@gmail.com



www.ijircce.com

Openmetadata and Datahub: A Comparative Evaluation of Open-Source Data Catalog Architectures for Automated Lineage, Discovery, and Governance in Modern Data Platforms

Venkata Vijay Satyanarayana Murthy Neelam

Senior Software Engineer - Cloud, Data, AI/ML & Generative AI, Atlanta, Georgia, USA

ABSTRACT: As organizations transition toward distributed data architectures, federated ownership models, and AI-augmented analytics, the selection of a data catalog platform has emerged as a foundational engineering decision with consequences spanning observability, compliance, query acceleration, and organizational data literacy. OpenMetadata and DataHub represent the two most architecturally mature open-source data catalog systems available as of 2024, each embodying distinct philosophical commitments about metadata ownership, lineage propagation, schema management, and extensibility. This paper delivers a comprehensive comparative evaluation across thirteen dimensions: architectural design, deployment complexity, metadata ingestion breadth, automated lineage capture, search and discovery capabilities, data quality integration, access control models, API surface and programmability, UI/UX characteristics, performance at scale, community ecosystem strength, operational observability, and enterprise governance readiness. Drawing on documented production deployments, community benchmarks, and direct architectural analysis of both platforms, we identify the conditions under which each system excels and formulate evidence-based selection criteria for engineering teams navigating this decision. Our evaluation concludes that OpenMetadata demonstrates superior out-of-the-box schema management, data quality embedding, and API-first design, while DataHub offers more mature graph-native lineage traversal, stronger push-based ingestion patterns, and a more established enterprise adoption base. Neither platform dominates unconditionally; team context, lineage fidelity requirements, and governance maturity are the decisive factors.

KEYWORDS: Data catalog, OpenMetadata, DataHub, data lineage, metadata management, data discovery, data governance, open-source, data mesh, schema registry, RBAC0

I. INTRODUCTION

The proliferation of cloud-native data infrastructure across organizations of every scale has created a metadata crisis of striking proportions. As data engineering teams provision data lakes in Amazon S3, Google Cloud Storage, and Azure Data Lake; orchestrate transformation pipelines through dbt, Apache Spark, and Airflow; serve analytical workloads via Snowflake, BigQuery, and Redshift; and expose real-time event streams through Apache Kafka and Confluent, the connective tissue that holds this complexity together - metadata - grows at a rate that dwarfs any individual team's capacity for manual curation.

Data catalog platforms address this metadata crisis by providing centralized registries that capture, store, relate, and surface information about data assets: what they contain, where they originate, who owns them, how they are transformed, what quality they exhibit, and who may access them. A mature data catalog converts opaque data silos into a discoverable, governed, and trusted data ecosystem.

1.1 The Stakes of Platform Selection

The choice between data catalog platforms is rarely reversible in the short term. Metadata ingestion pipelines, lineage integrations, access control policies, and organizational workflows become deeply entangled with the selected platform within twelve to eighteen months of deployment. An ill-suited selection imposes costs that extend well beyond licensing:

- Engineering teams spend weeks re-instrumenting source systems when lineage models prove incompatible with actual pipeline topologies
- Data stewards abandon catalog workflows when search relevance fails to return useful results for domain-specific terminology

- Compliance teams face audit failures when access control models cannot express the granularity required by GDPR, CCPA, or HIPAA
- Platform engineers find themselves maintaining unsupported custom connectors when the catalog vendor's connector ecosystem fails to cover essential source systems
- Executive stakeholders lose confidence in the catalog investment when query latency on lineage graphs degrades as metadata volume grows

OpenMetadata and DataHub both arose from the recognition that existing commercial catalog solutions impose prohibitive licensing costs while delivering mediocre extensibility, and that the open-source community was best positioned to build the connectivity and trust needed for a universal metadata standard to emerge.

1.2 Research Objectives and Scope

This evaluation pursues the following specific objectives:

- **Characterize the foundational architectural differences between OpenMetadata and DataHub and their downstream implications for scalability, extensibility, and operational complexity**
- **Evaluate the depth and reliability of automated data lineage capture across batch pipelines, streaming systems, SQL transformation tools, and BI platforms**
- **Assess search and discovery quality across large metadata corpora, including semantic search capabilities, faceted filtering, and glossary-driven navigation**
- **Compare data quality integration patterns and their impact on governance workflow automation**
- **Analyze API design philosophies and their implications for custom integrations and programmatic catalog management**
- **Formulate evidence-based platform selection criteria aligned with organizational maturity, team size, governance requirements, and technical stack characteristics**

This evaluation covers platform versions current as of Q4 2023 - OpenMetadata 1.2.x and DataHub 0.12.x - and draws on publicly documented production deployments, community forum analysis, GitHub issue trackers, and architectural documentation maintained by each project's core engineering teams.

II. BACKGROUND: THE METADATA MANAGEMENT LANDSCAPE

2.1 Evolution of Data Catalogs

Data catalog tooling has evolved through three recognizable generations. First-generation catalogs, exemplified by Apache Atlas (2015) and early commercial offerings, focused primarily on Hadoop ecosystem metadata capture. They stored table schemas, column definitions, and basic Hive partition information but lacked the connector breadth, search sophistication, and lineage depth required by modern polyglot data stacks.

Second-generation catalogs - including Alation (2012), Collibra (2008), and Informatica's Enterprise Data Catalog - brought richer business glossary management, collaborative annotation, and early lineage visualization. However, their commercial licensing models placed them out of reach for startups, mid-sized organizations, and teams committed to avoiding vendor lock-in. Their connector ecosystems, while broad, were updated on vendor-controlled schedules that frequently lagged behind rapidly evolving open-source tools.

Third-generation catalogs, of which OpenMetadata and DataHub are the most prominent open-source representatives, are distinguished by several architectural commitments:

- API-first design that treats every catalog operation as a first-class API call, enabling programmatic catalog management at scale
- Connector-as-code models where community members can contribute integrations following standardized connector interfaces
- Metadata standards orientation, with growing alignment toward the Open Metadata Initiative (OMI) and OpenLineage specification
- Embedded data quality, where quality metrics are not a separate concern but are stored and surfaced as metadata attributes on data assets
- Active metadata patterns, where the catalog acts not merely as a passive registry but as a trigger for automated governance actions, data contract enforcement, and pipeline orchestration

2.2 OpenLineage and the Standardization Imperative

A recurring pathology in earlier metadata ecosystems was the proliferation of proprietary lineage formats: each tool - Spark, Airflow, dbt, Flink - emitted lineage information in its own schema, requiring custom translation layers that

invariably fell behind as tools evolved. The OpenLineage specification [1], developed initially at Astronomer and contributed to the Linux Foundation in 2021, addresses this by defining a common event schema for lineage emission that any processing engine can implement.

Both OpenMetadata and DataHub have invested in OpenLineage compatibility, though their integration strategies differ significantly. DataHub's Marquez-compatible ingestion path consumes OpenLineage events through a dedicated Kafka topic, while OpenMetadata provides a native OpenLineage endpoint that maps incoming events directly to its Entity model. These design choices propagate into meaningful differences in lineage completeness, event ordering guarantees, and failure recovery behavior under high-throughput pipeline conditions.

HISTORICAL CONTEXT

LinkedIn's open-sourcing of DataHub in 2019 [2] and Collate's founding of OpenMetadata in 2021 [3] represent two distinct philosophical traditions: DataHub emerged from a web-scale social graph problem (modeling the metadata relationships across LinkedIn's thousands of datasets and microservices), while OpenMetadata was designed from inception as a universal open standard for metadata exchange, explicitly targeting the fragmented tooling landscape that had frustrated data engineering teams throughout the 2010s.

III. ARCHITECTURAL COMPARISON

3.1 OpenMetadata Architecture

OpenMetadata adopts a service-oriented architecture organized around a central Metadata Server that exposes a comprehensive REST API (OpenAPI 3.0-compliant) as its primary interface. All catalog operations - ingestion, discovery, lineage traversal, governance workflow execution - flow through this API surface, which serves as the contractual boundary between the catalog's core and every integration point.

Core Components

- Metadata Server: Stateless Java application (Spring Boot) that implements the full OpenMetadata API specification; horizontally scalable behind a load balancer
- MySQL or PostgreSQL: Primary relational storage for entity data, with an optimized schema that supports JSONB document storage for flexible entity attributes alongside normalized relational tables for efficient relational queries
- Elasticsearch / OpenSearch: Secondary search index populated via change-data-capture from the primary database; powers full-text search, faceted filtering, and relevance ranking
- Apache Airflow: Pipeline orchestration for scheduled metadata ingestion workflows; OpenMetadata ships a managed Airflow instance with pre-built connector DAGs
- Ingestion Framework: Python-based connector library (openmetadata-ingestion) that implements source connectors for 80+ data sources; connectors emit standardized Entity payloads to the Metadata Server API

Entity Model Philosophy

- OpenMetadata organizes metadata around a rich Entity taxonomy with approximately 35 distinct entity types as of version 1.2: Tables, Columns, Databases, Database Services, Pipelines, Pipeline Services, Topics, Charts, Dashboards, Dashboard Services, ML Models, ML Model Services, Containers, Data Products, Glossary Terms, Tags, Teams, Users, Policies, Roles, and more. Each entity type is defined by a JSON Schema specification that governs its attributes, relationships, and extension points. This schema-first approach means that any valid entity can be validated against its specification before ingestion, reducing the silent corruption that plagues schema-lax metadata stores.

KEY FINDING

OpenMetadata's entity model enforces schema validation at ingestion time, catching malformed metadata before it contaminates the catalog. This design choice trades some ingestion flexibility for consistency guarantees that become critically important at scale, where a single malformed lineage edge can corrupt downstream impact analysis for thousands of downstream assets.

3.2 DataHub Architecture

DataHub's architecture reflects its LinkedIn origins as a distributed system designed for horizontal scalability from day one. Rather than a single metadata server, DataHub decomposes metadata management into a set of discrete microservices that communicate primarily through Apache Kafka, establishing an event-driven foundation that enables both real-time metadata updates and historical event replay.

Core Components

- **GMS (Generalized Metadata Service):** The central metadata storage and query service; exposes a REST API and GraphQL endpoint for entity CRUD operations and relationship traversal
- **MCE (Metadata Change Event) Kafka Topic:** The primary ingestion channel; producers emit metadata change proposals (MCPs) as Avro-encoded Kafka messages; GMS consumes and applies them
- **MAE (Metadata Audit Event) Kafka Topic:** Downstream consumers of committed metadata changes; powers secondary indices, notifications, and external system synchronization
- **Elasticsearch:** Full-text search index; populated from MAE consumers; supports entity search, faceted filtering, and DataHub's "Browse" tree navigation
- **Neo4j or graph-compatible backend:** DataHub 0.10+ supports a pluggable graph backend; the default uses Elasticsearch-based graph queries, with Neo4j available for deployments requiring deep graph traversal performance
- **datahub-frontend:** React-based single-page application; proxies API calls to GMS through a dedicated Play Framework service
- **datahub-actions:** Event-driven automation framework that subscribes to MAE events and executes configured actions (Slack notifications, JIRA ticket creation, policy enforcement) in response to metadata changes

Metadata Aspect Model

DataHub organizes metadata through the concept of "Aspects" - discrete units of metadata that attach to an Entity. An Entity (such as a Dataset or Data Pipeline) may have many Aspects: SchemaMetadata, DatasetProperties, Ownership, GlobalTags, Glossary Terms, Lineage, Status, and custom aspects defined by users. This aspect-oriented model provides several architectural advantages:

- Aspects can be updated independently without affecting other metadata attributes of the same entity
- The Kafka-based ingestion model allows aspects from different source systems to be merged into a unified entity view
- Custom aspects enable teams to extend the metadata model without forking the core codebase
- Aspect versioning enables point-in-time reconstruction of entity state - effectively providing time-travel metadata queries

3.3 Architectural Trade-off Summary

The architectural divergence between the two platforms creates a set of fundamental trade-offs that engineering teams must evaluate against their specific operational context:

- **Deployment complexity:** DataHub's microservices architecture requires orchestrating Kafka, Elasticsearch, GMS, the frontend, and optionally Neo4j as distinct services with independent scaling, health checking, and configuration management. OpenMetadata's more consolidated deployment (Metadata Server + MySQL/Postgres + Elasticsearch + Airflow) presents a simpler operational surface area, though it sacrifices the fine-grained horizontal scaling that DataHub's model enables
- **Ingestion coupling:** OpenMetadata's synchronous REST API ingestion provides immediate consistency - entities are queryable the moment ingestion completes - while DataHub's async Kafka-based ingestion delivers higher throughput at the cost of ingestion lag; datasets ingested via MCE may not appear in search for seconds to minutes after emission depending on consumer lag
- **Schema governance:** OpenMetadata's JSON Schema-validated entity model enforces structural correctness at ingestion time; DataHub's Avro-schema-governed aspects enforce correctness at the Kafka producer level but permit more flexible schema evolution through Avro's compatibility modes
- **State management:** DataHub maintains a full event log in Kafka, enabling complete metadata history reconstruction and consumer replay; OpenMetadata stores entity history in its relational database through entity change events, providing similar time-travel capabilities but without the distributed log's partition tolerance properties

IV. METADATA INGESTION AND CONNECTOR ECOSYSTEMS

4.1 OpenMetadata Connector Framework

OpenMetadata's ingestion framework implements connectors as Python classes that extend a Source base class, yielding Entity payloads to a configurable Sink (typically the REST API sink that delivers payloads to the Metadata Server). This design enables connectors to be developed, tested, and contributed without any changes to the core server, accelerating community-driven connector development.

Supported Source Categories (OpenMetadata 1.2.x)

- Database Services: MySQL, PostgreSQL, Snowflake, BigQuery, Redshift, Databricks SQL, SQL Server, Oracle, MongoDB, DynamoDB, Cassandra, Presto, Trino, Hive, Athena, MariaDB, SQLite, ClickHouse, Druid, Greenplum, Impala, SingleStore, and more (62 database connectors as of v1.2)
- Dashboard Services: Apache Superset, Metabase, Tableau, PowerBI, Redash, QuickSight, Looker, Mode, Lightdash, Domo - with lineage from dashboard queries back to source tables
- Pipeline Services: Apache Airflow (native plugin captures DAG-level and task-level lineage at run time), dbt (parses manifest.json and catalog.json for model lineage), Prefect, Dagster, Apache Spark (through OpenLineage events), Apache Flink, Glue, Fivetran
- Messaging Services: Apache Kafka, Confluent Schema Registry (captures topic schemas and subject versions), Redpanda, Kinesis
- ML Platforms: MLflow (captures model versions, parameters, metrics, and feature lineage), SageMaker, Vertex AI, custom REST endpoints
- Storage Services: Amazon S3, Google Cloud Storage, Azure Data Lake Storage - with container/object metadata capture
- Metadata Services: Cross-catalog ingestion from Amundsen, Atlas, and other catalog systems to enable migration or federation scenarios

Ingestion Workflow Patterns

- Workflow-based scheduling: Each connector configuration produces an Airflow DAG that runs on a configurable schedule (hourly, daily, weekly); the OpenMetadata-managed Airflow instance handles DAG registration, execution, and failure alerting
- Incremental ingestion: Connectors support both full-refresh and incremental ingestion modes; incremental mode captures only schema changes, ownership updates, and tag modifications since the last successful run, reducing API call volume and ingestion runtime for large source systems
- Metadata versioning: Each ingestion run produces a versioned snapshot of entity state; the Metadata Server maintains entity change history, enabling diff-based change detection and audit trail generation

4.2 DataHub Ingestion Framework

DataHub's ingestion framework (datahub-ingestion Python library) follows a similar source-transformer-sink pipeline model but is oriented around the MCP (Metadata Change Proposal) as the atomic unit of metadata change. Connectors emit MCPs to either the REST sink (synchronous) or the Kafka sink (asynchronous, higher throughput), giving teams the flexibility to trade consistency for throughput based on their use case.

Supported Source Categories (DataHub 0.12.x)

- Database & Warehouse Sources: Snowflake (with query-log-based lineage), BigQuery (with usage statistics and lineage from INFORMATION_SCHEMA), Redshift, Databricks, SQL Server, PostgreSQL, MySQL, Oracle, Hive, Trino, Presto, Athena, Glue Data Catalog, and 55+ additional database sources
- BI & Visualization Platforms: Tableau (deep lineage to underlying datasets via Tableau Metadata API), PowerBI (lineage to source datasets via PowerBI Scanner API), Looker (LookML model lineage), Superset, Mode, Metabase, Sigma
- Transformation & Orchestration: dbt (native support for manifest.json, catalog.json, and sources.json; captures model-to-model lineage and test results), Airflow (lineage from task inlets/outlets), Spark (via OpenLineage listener)
- Streaming Platforms: Apache Kafka (topic schemas from Schema Registry; consumer group offset tracking), Pulsar, Kinesis
- Data Lakes: S3 (path-based dataset discovery), Azure Data Lake, Google Cloud Storage, Delta Lake, Apache Iceberg (table-level metadata from catalog)
- Custom Sources: DataHub's Python SDK enables custom source development in approximately 200 lines of code for well-structured APIs; the recipe-based configuration system makes custom sources deployable without server-side changes

Push-Based Ingestion Advantages

DataHub's architectural emphasis on push-based metadata emission (via the datahub Python client or OpenLineage Kafka events) enables a fundamentally different lineage capture model than pull-based connector polling. When Spark jobs, dbt runs, Airflow tasks, and Kafka consumers emit lineage events in real time, the catalog reflects pipeline state within seconds of execution rather than waiting for the next scheduled ingestion run. This real-time metadata push model delivers several meaningful advantages:

- Lineage freshness: Column-level lineage is captured at job execution time from the actual query plan, not inferred from SQL parsing after the fact - eliminating a category of lineage inaccuracy that plagues post-hoc SQL parsing approaches
- Failure attribution: When a downstream dataset fails SLA, DataHub's real-time lineage reflects which upstream job was executing (and potentially failing) at the time, enabling faster root-cause attribution
- Schema drift detection: As jobs execute, emitted schema aspects reflect the actual schemas encountered at runtime, catching schema drift between expected and actual schemas in near-real time

KEY FINDING

DataHub's push-based, event-driven ingestion model captures lineage at job execution time with sub-minute freshness, while OpenMetadata's pull-based scheduled ingestion provides more predictable operational behavior but at the cost of lineage staleness - typically measured in hours for daily-scheduled connectors.

V. AUTOMATED DATA LINEAGE: DEPTH, FIDELITY, AND TRAVERSAL

5.1 Lineage Granularity Levels

Data lineage in modern data platforms must be captured and represented at multiple granularity levels simultaneously, as different stakeholders require different resolution:

- Table/Dataset-level lineage: Answers "which tables flow into which other tables?" - sufficient for high-level impact analysis and dependency documentation, but insufficient for column-level impact and data masking decisions
- Column-level lineage: Answers "which source columns influence which target columns?" - required for GDPR right-to-erasure propagation, data masking compliance, and fine-grained impact analysis when source schemas evolve
- Run-level lineage: Answers "which specific pipeline execution produced this version of the dataset?" - required for debugging, SLA attribution, and data quality incident investigation
- Field-transform lineage: Answers "what transformations were applied to a column in transit?" - the most granular and most valuable for compliance, though the most difficult to capture automatically

5.2 OpenMetadata Lineage Capabilities

OpenMetadata captures lineage through multiple complementary pathways, each optimized for different source system characteristics

SQL Query-Based Lineage Parsing

For SQL-native transformation environments, OpenMetadata's query lineage parser analyzes SQL query text to extract table and column-level lineage. The parser handles complex SQL constructs including CTEs, subqueries, window functions, UNION operators, and conditional expressions. Supported databases expose query history through system views (Snowflake's `QUERY_HISTORY`, BigQuery's `INFORMATION_SCHEMA.JOBS`, Redshift's `STL_QUERY`), which the ingestion connector harvests and feeds to the lineage parser.

- Supported SQL dialects: Snowflake SQL, BigQuery Standard SQL, Redshift SQL (PostgreSQL-compatible), Spark SQL, Hive QL, Trino/Presto SQL, ANSI SQL
- Column-level lineage accuracy: For straightforward `SELECT-FROM-WHERE` patterns, OpenMetadata achieves near-100% accuracy; complex expressions involving user-defined functions, dynamic SQL, or multi-level CTEs require manual supplementation
- Query coverage: OpenMetadata ingests the last N days of query history (configurable; default 1 day for incremental runs); for systems without query history APIs, lineage relies on dbt manifest files or manually specified edges

dbt Native Lineage

OpenMetadata's dbt connector is one of the most complete dbt integrations in the catalog ecosystem, ingesting metadata from dbt's `manifest.json`, `catalog.json`, `sources.json`, and `run_results.json` artifacts:

- Model-level lineage: Complete DAG of dbt model dependencies, including ephemeral models, seeds, and source references
- Column-level lineage: Extracted from compiled SQL in the manifest; limited by dbt's own column resolution, which may not traverse through complex `ref()` expressions
- Test results: dbt test pass/fail outcomes ingested as data quality dimensions on target tables
- Descriptions and documentation: dbt model and column descriptions propagated to OpenMetadata entity descriptions, populating the catalog with developer-authored documentation

5.3 DataHub Lineage Capabilities

DataHub's lineage model is built on a directed acyclic graph (DAG) where Datasets, Data Jobs, Data Flows, and Charts are the primary node types, and lineage edges carry directional semantics (UPSTREAM, DOWNSTREAM) with optional transformation context. The graph representation enables multi-hop traversal queries - "show me all datasets that are downstream of this source table, up to 5 hops" - that are foundational to impact analysis at enterprise scale.

Graph-Native Lineage Traversal

- Multi-hop lineage: DataHub's lineage API supports configurable depth traversal with cycle detection, returning the full upstream or downstream lineage DAG as a structured response
- Cross-platform lineage: A single lineage path can traverse Kafka topics, Spark jobs, Delta Lake tables, dbt models, and Tableau dashboards - each represented as typed entities with platform-specific metadata
- Column-level lineage: Supported through the UpstreamLineage aspect's fineGrainedLineage field; column pairs are stored as explicit (sourceField, destinationField) tuples attached to the transformation job entity
- Lineage visualization: DataHub's Impact Analysis UI renders the lineage DAG as an interactive graph with collapsible depth levels, entity type filtering, and direct navigation to individual entity pages

Real-Time Lineage via OpenLineage

- Kafka-based event ingestion: Spark jobs instrumented with the datahub-spark-agent emit OpenLineage-compatible events to a Kafka topic; the DataHub OpenLineage consumer translates these to MCPs and applies them to the graph in near-real time
- Airflow integration: DataHub's Airflow plugin captures task-level lineage from Airflow's XCOM and inlets/outlets metadata, creating DataJob entities for each task and linking them to upstream/downstream Dataset entities
- dbt integration: DataHub's dbt source ingests manifest artifacts and maps dbt model dependencies to DataHub lineage edges; column-level lineage from compiled SQL is extracted by a dedicated SQL parsing step

LINEAGE COMPLETENESS NOTE

Both platforms face the same fundamental challenge: lineage completeness is bounded by instrumentation coverage. A Spark job that reads from S3 without emitting OpenLineage events, or a Python script that reads from Snowflake via JDBC without query logging, will produce no automatic lineage. Teams should budget for instrumentation effort proportional to pipeline diversity and should treat automated lineage as a starting point that manual curation supplements, not a complete solution.

VI. SEARCH AND DISCOVERY CAPABILITIES

6.1 OpenMetadata Search Architecture

OpenMetadata's search capability is built on Elasticsearch/OpenSearch and configured with a sophisticated mapping schema that indexes entity metadata across multiple fields with different analyzers, enabling both exact-match and fuzzy full-text search. The search system is designed around the concept of "discovery journeys" - structured pathways through which data consumers navigate from a business concept to a specific data asset.

Search Features

- Full-text search across entity names, descriptions, column names, column descriptions, glossary term definitions, and tag labels with relevance ranking weighted toward entity name matches
- Faceted filtering by entity type (Table, Dashboard, Topic, Pipeline, ML Model), service/platform, tier (data quality tier), tags, owners, and domain - enabling progressive narrowing of large result sets
- Glossary-driven discovery: Business glossary terms are linked to data assets, enabling non-technical users to search by business concept ("customer lifetime value") and navigate to the underlying tables and dashboards
- Data Insights dashboard: Aggregated analytics showing catalog coverage (percentage of tables with descriptions, owners, tags), most-viewed assets, most-queried tables (via usage statistics), and team-level catalog activity
- Explore by domain: Hierarchical domain structure organizes assets into business domains (Marketing, Finance, Operations), enabling domain-centric browsing that aligns with data mesh organizational patterns
- Recent and popular: Usage-based ranking surfaces frequently queried and recently accessed assets as default-ranked discovery starting points
- Column-level search: Search terms matched against column names and descriptions enable discovery of assets by their attributes, not just their entity-level names - critical when the desired table name is unknown but a column name is recognizable

6.2 DataHub Search Architecture

DataHub's search is also Elasticsearch-backed but is distinguished by a more sophisticated query rewriting layer (DataHub Query Builder) that translates user queries into multi-field Elasticsearch queries with platform-specific boosting, entity type weighting, and ownership-based personalization. The search infrastructure has been battle-tested at LinkedIn scale and includes several features reflecting that heritage:

Search Features

- Union-find based entity deduplication: When the same logical dataset exists under multiple aliases or is ingested from multiple source systems, DataHub's entity key resolution prevents duplicate entries from polluting search results
- Autocomplete with entity type hints: As users type, autocomplete suggestions are categorized by entity type, reducing ambiguity when the same term applies to a table, a column, and a glossary term simultaneously
- Search by URN: DataHub's Uniform Resource Name scheme (urn:li:dataset:(urn:li:dataPlatform:snowflake,mydb.schema.mytable,PROD)) enables precise addressing of specific entities, supporting programmatic access and cross-system linking
- Filter combinations: DataHub's filter sidebar supports AND/OR logic across filter dimensions, enabling complex filter expressions not available in OpenMetadata's simpler faceted interface
- GraphQL search API: DataHub exposes its full search capability through GraphQL, enabling frontend applications and external tools to compose complex entity queries with relationship traversal in a single request
- Personalized feed: DataHub's home feed surfaces entities related to the user's ownership assignments, subscribed domains, and recently viewed assets - providing a personalized starting point for daily catalog engagement
- Business glossary integration: Glossary Terms appear as searchable entities; assets tagged with glossary terms surface in glossary term search results, enabling bidirectional navigation between business concepts and data assets

KEY FINDING

OpenMetadata's search excels at description-rich environments where data assets have been documented with verbose, searchable text. DataHub's search outperforms in high-entity-count environments (100,000+ entities) where the union-find deduplication and sophisticated query rewriting maintain relevance that degrades in OpenMetadata's simpler relevance model.

VII. DATA GOVERNANCE: POLICIES, RBAC, AND COMPLIANCE

7.1 OpenMetadata Governance Model

OpenMetadata implements governance through a hierarchical policy and role model that encompasses Teams, Users, Roles, Policies, and Rules. This model is designed to support both coarse-grained organizational access control and fine-grained asset-level permission management within a unified policy framework.

Access Control Components

- Roles: Named permission sets (e.g., "Data Consumer," "Data Steward," "Data Owner") that aggregate a collection of Policies; roles are assigned to Users or Teams
- Policies: Rule collections that govern which operations a role may perform on which entity types; operations include ViewAll, EditDescription, EditOwners, EditTags, EditGlossaryTerms, EditLineage, EditQueries, and EditCustomFields
- Rules: Atomic permission expressions combining Resource (entity type or specific entity), Operation, and Effect (Allow/Deny); deny rules take precedence over allow rules, enabling fine-grained exception management
- Teams hierarchy: Organizations model their structure as a tree of Teams, with Roles assignable at any level; inherited permissions propagate down the team hierarchy, enabling organizational-scale permission management without per-user configuration
- Domain-level policies: Governance policies can be scoped to specific Domains, implementing the data mesh principle that domain owners control access to their own data products
- Data Tier classification: Assets are classified into Tiers (Tier 1 through Tier 5, or custom tiers) indicating data quality and business criticality; tier-based policies enable routing of high-tier assets to stricter review and access workflows

Governance Workflows

- Tag-based policy enforcement: Sensitivity tags (PII, Confidential, Internal) attached to tables and columns serve as policy anchors; governance workflows monitor tag application and trigger access review notifications
- Glossary term approval workflow: New glossary term proposals enter a multi-stage review workflow (Proposal → Review → Approved), ensuring that business vocabulary is curated and consensus-driven

- Data product ownership: The Data Product entity type formalizes ownership boundaries, enabling domain teams to declare which datasets constitute their data product and what quality and SLA commitments they make to consumers
- Alerts and notifications: OpenMetadata's activity feeds and email/Slack notification integrations surface governance events (ownership changes, tag applications, description edits) to configured stakeholders

7.2 DataHub Governance Model

DataHub's governance model is built around the concept of "Metadata Policies" - server-side rules that govern what operations specific users or groups may perform on specific entity types or individual entities. DataHub's governance capabilities have been expanded significantly in the 0.10-0.12 release cycle, reflecting enterprise customer feedback about compliance requirements.

Access Control Components

- Platform Policies: Control access to platform-level operations (manage users, manage groups, manage policies, manage secrets, create domains) - governing who can administer the catalog itself
- Metadata Policies: Control access to entity-level operations (view entity, edit entity properties, edit tags, edit owners, edit lineage, manage assertions); policies reference entity type filters and specific entity URN filters
- Actors: Policies apply to Users, Groups, or the special ALL_USERS actor; Group membership is synchronized from LDAP/Active Directory or SCIM-compatible IdP integrations
- Privileges: Fine-grained operation set covering VIEW_ENTITY_PAGE, EDIT_ENTITY_TAGS, EDIT_ENTITY_GLOSSARY_TERMS, EDIT_ENTITY_OWNERS, EDIT_ENTITY_DOCS, EDIT_ENTITY_STATUS, PROPOSE_ENTITY_TAGS, PROPOSE_GLOSSARY_TERMS - distinguishing between write operations and proposal operations that trigger human review
- Data Contracts: DataHub's Data Contracts feature (beta in 0.12) formalizes schema, freshness, and volume expectations for datasets as machine-readable contracts; contract violations trigger policy-configured actions through the datahub-actions framework

Governance Automation via datahub-actions

- Event-driven policy enforcement: datahub-actions subscribes to MetadataChangeLog events from Kafka and executes configured Actions (Python callables) in response; pre-built actions include Slack notifications, JIRA ticket creation, and Teams alerts
- Propagation actions: Custom actions can propagate tags, owners, or terms from upstream entities to downstream entities along lineage edges - implementing policy-driven metadata inheritance
- Assertion monitoring: DataHub's Assertions framework defines freshness, volume, column, and SQL-based assertions on datasets; assertion results are stored as metadata aspects and surfaced in the entity UI and governance dashboards
- Incident management: When assertion failures exceed configured thresholds, DataHub can auto-create Incidents - tracked governance events with assignee, severity, and resolution status - linking the incident to the affected dataset for full audit trail

VIII. DATA QUALITY INTEGRATION

8.1 OpenMetadata Native Data Profiling

OpenMetadata has invested heavily in native data quality and profiling capabilities that are embedded directly in the catalog, rather than requiring integration with an external quality platform. This design philosophy - that data quality metadata belongs in the catalog, not siloed in a separate quality system - drives OpenMetadata's most differentiating capability cluster.

Profiling Capabilities

- Column profiling: Automated statistical profiling computing row count, null count, null percentage, unique count, distinct count, min, max, mean, median, stddev, and quartiles for every column in a table on each ingestion run
- Custom metrics: User-defined SQL expressions evaluated during profiling runs, storing the computed value as a metric time series on the table entity; enables business-rule metrics ("active accounts percentage") alongside standard statistical profiles
- Data quality tests: SQL-based and Python-based test definitions attached to tables and columns; test results (pass/fail/aborted) stored as DQ test run entities with time-series tracking
- Test suites: Logical groupings of related tests that can be executed as atomic units; test suite results aggregate to a table-level quality score

- Quality dashboards: Built-in dashboards visualizing test pass rates, profiling coverage, column-level quality trends, and team-level quality scores - enabling data quality as a first-class observable dimension of catalog health
- dbt test integration: dbt test results (unique, not_null, accepted_values, relationships, and custom generic tests) imported through the dbt connector populate OpenMetadata's quality framework, unifying dbt-managed and catalog-native quality results under a single interface

8.2 DataHub Data Quality Integration

DataHub approaches data quality through its Assertions framework and external integrations rather than native profiling execution. While OpenMetadata executes quality checks within the catalog infrastructure, DataHub positions itself as the aggregation layer for quality results produced by purpose-built quality systems.

Assertions Framework

- Freshness assertions: Define expected update frequency for datasets (e.g., "this table should be updated at least once every 24 hours"); freshness is evaluated against the table's last modification timestamp from the source system
- Volume assertions: Define expected row count ranges (minimum, maximum, or growth rate); evaluated against profiling results or source system statistics
- Column assertions: Define column-level constraints (null rate below threshold, value range, pattern matching); evaluated by executing SQL checks against the source system
- Custom SQL assertions: Arbitrary SQL expressions evaluated against the source system, with pass/fail determined by whether the result meets a configured condition
- External assertion integration: Great Expectations, dbt tests, and Soda Core can push their test results to DataHub as assertion results through the datahub Python client, populating the assertion framework without requiring DataHub to execute the checks itself

DESIGN PHILOSOPHY CONTRAST

OpenMetadata's native quality execution model means that a team using OpenMetadata as their sole data infrastructure platform can implement a complete data quality workflow without adopting additional tools. DataHub's integration-focused approach is better suited to organizations that have already invested in dedicated quality platforms (Great Expectations, Monte Carlo, Soda) and want the catalog to serve as the unified observability surface rather than the quality execution engine.

IX.API DESIGN, PROGRAMMABILITY, AND EXTENSIBILITY

9.1 OpenMetadata API Surface

OpenMetadata's API-first design philosophy manifests in one of the most complete and consistently designed REST API surfaces in the data catalog ecosystem. The OpenAPI 3.0 specification that governs every endpoint is generated from Java code annotations, ensuring that the API specification is always synchronized with the server implementation.

- Entity CRUD APIs: Standardized create/read/update/delete/patch endpoints for all 35+ entity types following a consistent URL schema (/api/v1/{entityType}); every entity type exposes the same CRUD surface, minimizing the learning curve for API consumers
- Search API: Parameterized search endpoint supporting entity type filtering, field-level querying, sort ordering, and pagination; returns entity summaries with configurable field projection
- Lineage API: Directed lineage traversal API returning upstream/downstream entity graphs with configurable depth and entity type filtering; supports both column-level and table-level lineage queries
- Ingestion API: Programmatic creation, configuration, triggering, and monitoring of ingestion pipelines through the API; enables CI/CD-driven catalog management where pipeline definitions are version-controlled and deployed as code
- Events API: Server-Sent Events (SSE) endpoint exposing entity change events in real time; enables event-driven integrations that react to catalog changes without polling
- Python SDK: openmetadata-sdk Python package provides strongly typed client classes for all entity types, reducing boilerplate for Python-native integrations
- Java client: Java SDK available for JVM-based integration scenarios

9.2 DataHub API Surface

DataHub exposes three distinct API surfaces, each optimized for different consumption patterns:

REST API (GMS)

- Entity read/write through generic entity endpoints and entity-specific endpoints; the generic /entities endpoint accepts any valid Entity + Aspect combination, enabling a single code path for all entity types
- Lineage queries through dedicated /lineage endpoint with depth, direction, and entity type filter parameters
- Search through /entities/search endpoint returning ranked results with facet counts

GraphQL API

- Schema-driven query composition enabling clients to fetch exactly the entity fields they need - including relationship traversal - in a single round trip; eliminates the over-fetching that REST APIs exhibit when clients need a subset of entity attributes
- Mutation support for entity updates, tag applications, owner assignments, and glossary term linking - providing a GraphQL-native alternative to REST mutations
- Subscription support (in preview) enabling real-time entity change notifications through GraphQL subscriptions over WebSockets
- DataHub's React frontend is itself a GraphQL consumer, ensuring that the GraphQL API covers the full feature set exposed by the UI rather than lagging behind

Python Ingestion SDK

- datahub Python library providing MetadataChangeProposalWrapper and specialized builder classes for each entity/aspect combination
- Recipe-based CLI (datahub ingest -c recipe.yaml) for declarative ingestion configuration without writing Python code
- Pipeline-style ingestion with transformer support enabling metadata enrichment (adding tags, owners, terms based on entity attributes) as a processing step between source and sink

X. DEPLOYMENT COMPLEXITY AND OPERATIONAL CONSIDERATIONS

10.1 OpenMetadata Deployment

OpenMetadata provides first-class Kubernetes deployment support through a comprehensive Helm chart (open-metadata/openmetadata) that packages the Metadata Server, MySQL, Elasticsearch, and the managed Airflow instance as a coordinated release. For evaluation and smaller deployments, a Docker Compose configuration packages all components for single-machine deployment with persistent volume configuration.

Deployment Options

- Docker Compose: Single-command deployment for development and evaluation; suitable for teams ingesting fewer than 50 data sources with moderate catalog scale (under 100,000 entities)
- Helm/Kubernetes: Production deployment with configurable resource requests, horizontal pod autoscaling for the Metadata Server, and external MySQL/Postgres and Elasticsearch cluster support; the managed Airflow instance is recommended for most deployments, with external Airflow support for teams with existing Airflow infrastructure
- Managed deployment: Collate (the commercial entity behind OpenMetadata) offers a fully managed SaaS version that eliminates infrastructure management for teams that prefer operational simplicity over self-hosting control
- OpenMetadata Sandbox: A public sandbox environment maintained by the OpenMetadata team enables evaluation without any local deployment, accelerating proof-of-concept timelines

Operational Considerations

- Database sizing: MySQL/PostgreSQL is the operational bottleneck for large catalogs; teams ingesting 500,000+ entities should provision dedicated database instances with at least 16 GB RAM and SSD storage; the entity JSON storage pattern enables flexible schema but increases storage volume relative to fully normalized schemas
- Elasticsearch cluster sizing: A two-node Elasticsearch cluster (2 shards, 1 replica) serves most deployments up to 500,000 entities; beyond this scale, dedicated Elasticsearch clusters with hot-warm tier architectures are recommended
- Airflow resource allocation: Each concurrent ingestion workflow consumes 500MB-2GB of memory depending on source system size; teams with many concurrent daily ingestion workflows should scale the Airflow worker pool accordingly
- Upgrade path: OpenMetadata maintains migration scripts for each minor version upgrade; the upgrade process involves database schema migration (typically 5-30 minutes for large catalogs) followed by rolling Metadata Server pod replacement

10.2 DataHub Deployment

DataHub's microservices architecture creates a more complex deployment footprint that reflects its design heritage as a LinkedIn-scale system. Production deployments require coordinating Kafka, Zookeeper (or KRaft for Kafka 3.x), Elasticsearch, MySQL or PostgreSQL, the GMS service, the frontend service, the actions service, and optionally Neo4j.

Deployment Options

- Docker Compose (quickstart): datahub docker quickstart command uses a pre-configured Docker Compose stack for evaluation; the quickstart configuration is not recommended for production due to ephemeral storage and single-node Kafka configuration
- Helm/Kubernetes: Comprehensive Helm charts in the datahub-helm repository; each microservice is packaged as a separate chart with configurable resource allocation, enabling precise scaling of individual components; the Kafka deployment is the most operationally complex component, requiring careful topic configuration, replication factor settings, and consumer group management
- Amazon Managed Streaming for Kafka (MSK): AWS deployments frequently use MSK as the Kafka provider, eliminating Kafka/Zookeeper operational burden at the cost of AWS vendor dependency
- Acryl Data: The commercial entity behind DataHub offers a managed cloud deployment (Acryl Cloud) providing the full DataHub feature set plus enterprise additions without self-hosting complexity

Operational Considerations

- Kafka topic management: DataHub creates approximately 10 Kafka topics at initialization; these topics must be pre-created with appropriate partition counts and replication factors before GMS startup; incorrect partition counts are difficult to change without reprocessing the full event log
- Consumer lag monitoring: The lag of the MAE consumer group is a critical operational metric; sustained consumer lag indicates that metadata changes are accumulating in Kafka faster than GMS can process them, leading to search index staleness and lineage display delays
- GMS resource allocation: GMS is the most resource-intensive DataHub component; production deployments typically allocate 8-16 GB heap to GMS with a corresponding container memory limit of 12-24 GB; heap exhaustion under high ingestion load is a common source of production incidents
- Schema Registry dependency: DataHub's Avro-encoded Kafka messages require a Schema Registry (Confluent Schema Registry or AWS Glue Schema Registry) for serialization/deserialization; schema compatibility settings must be configured to FULL_TRANSITIVE to prevent producer-consumer schema incompatibilities after version upgrades

XI. COMMUNITY ECOSYSTEM AND ENTERPRISE ADOPTION

11.1 OpenMetadata Community

- GitHub repository (open-metadata/OpenMetadata): 4,200+ stars as of Q4 2023; active issue tracker with typical P1 bug resolution within 72 hours by core team
- Slack community: 4,000+ members; active #support and #connectors channels with community response times typically under 4 hours during business hours
- Connector contributions: 40% of the 82 connectors in OpenMetadata 1.2 were contributed by community members rather than the Collate core team - evidence of healthy contributor engagement
- Release cadence: Monthly minor releases with biweekly patch releases; the 1.x release cycle has maintained semantic versioning with documented breaking change policies
- Documentation quality: OpenMetadata's documentation site (docs.open-metadata.org) covers every connector with configuration reference, deployment guides, and architecture documentation; video tutorials and blog posts supplement written documentation
- Enterprise adoption: Documented production deployments at Salesforce, Netflix (evaluation), Booking.com, and numerous fintech and healthtech companies; catalog scale ranges from 10,000 to 2,000,000+ entities across documented deployments

11.2 DataHub Community

- GitHub repository (datahub-project/datahub): 8,900+ stars as of Q4 2023 - the largest GitHub star count of any open-source data catalog project; reflects DataHub's longer history and LinkedIn brand association
- Slack community: 7,000+ members across the DataHub Slack workspace; dedicated channels for #ingestion, #frontend, #backend, #deployment, and vertical-specific channels
- Connector ecosystem: 60+ source connectors; the datahub-contrib organization hosts community-contributed connectors that are vetted but not maintained by the Acryl core team
- Enterprise adoption: Production deployments at LinkedIn (origin; 1M+ entities), Visa, Optum, Peloton, Saxo Bank, Hurb, and numerous global financial institutions; DataHub's longer history gives it a broader set of documented enterprise reference implementations
- DataHub Town Halls: Monthly virtual community events featuring roadmap previews, community Q&A, and production user presentations - fostering community engagement beyond GitHub issue resolution

- OpenLineage Working Group participation: DataHub contributors are active participants in the OpenLineage specification working group, contributing to the standard that both platforms implement

XII. PLATFORM SELECTION FRAMEWORK

12.1 Choose OpenMetadata When...

- Your team prioritizes an API-first catalog where every operation is a documented REST call and programmatic catalog management is a core workflow - OpenMetadata's OpenAPI specification and Python/Java SDKs are the most complete in the ecosystem
- Native data quality execution is a requirement and you prefer not to adopt a separate quality platform; OpenMetadata's embedded profiling, test suites, and quality dashboards deliver a complete quality observability workflow within the catalog
- Your data stack is SQL-heavy and dbt-centric; OpenMetadata's dbt integration captures the richest set of dbt metadata artifacts and propagates them to the most catalog dimensions
- Your organization is implementing a data mesh architecture with formal Domain ownership boundaries; OpenMetadata's Data Product entity type and domain-scoped governance policies map directly to data mesh principles
- Deployment simplicity is important and your team does not have dedicated infrastructure engineers for Kafka operations; OpenMetadata's more consolidated architecture reduces operational surface area
- You require a rich schema versioning and schema change tracking capability; OpenMetadata's entity history and schema diff visualizations provide granular visibility into schema evolution over time
- Your evaluation timeline is short and you need a working catalog quickly; OpenMetadata's sandbox environment and Docker Compose quickstart enable a functional evaluation within hours

12.2 Choose DataHub When...

- Lineage graph depth and traversal performance are primary requirements; DataHub's graph-native architecture and multi-hop lineage API outperform OpenMetadata at high entity counts and deep lineage graphs
- Real-time, push-based metadata emission is architecturally preferable to scheduled pull ingestion; DataHub's Kafka-based event stream enables sub-minute metadata freshness from instrumented pipelines
- Your organization has already adopted Kafka as a core infrastructure component and has the operational expertise to manage a Kafka-based microservices deployment
- You require a GraphQL API for flexible entity queries with relationship traversal; DataHub's GraphQL surface is more mature than OpenMetadata's REST-only interface for complex relational queries
- You are building a large-scale catalog with 500,000+ entities and require fine-grained horizontal scaling of individual catalog components - DataHub's microservices decomposition enables component-level scaling that OpenMetadata's monolithic server cannot match
- Your team wants to implement event-driven governance automation where metadata changes trigger automated downstream actions; DataHub's datahub-actions framework is more mature than OpenMetadata's notification system for complex automation workflows
- Enterprise reference implementations and a large community of peer adopters are important evaluation criteria; DataHub's broader enterprise adoption base and 4+ year production history at LinkedIn provide more battle-tested reference implementations

12.3 Hybrid and Migration Considerations

Some organizations have successfully deployed both catalogs in a federated model - DataHub as the enterprise lineage graph serving the central data platform team, and OpenMetadata as the domain-level catalog for individual business unit teams who prefer its API-first integration model and embedded quality features. This approach is architecturally supported by both platforms' cross-catalog ingestion capabilities, though it introduces synchronization complexity and potential consistency challenges between catalog states.

Teams considering migration between platforms should evaluate the following migration dimensions:

- Entity metadata portability: Both platforms expose entity metadata via REST APIs, enabling extraction of descriptions, owners, tags, and glossary terms for programmatic re-ingestion into the target platform
- Lineage portability: Lineage graph export is less standardized; teams should budget for re-ingestion from source systems rather than attempting direct lineage graph migration between catalog formats
- Custom field migration: Custom properties and custom entity types defined in the source catalog require mapping to the target platform's extension mechanisms; this is typically the most labor-intensive migration component
- Governance policy migration: Access control policies, role definitions, and team structures must be re-expressed in the target platform's policy model; direct policy migration is not supported by either platform

XIII. FUTURE DIRECTIONS AND RESEARCH GAPS

13.1 Emerging Capabilities

Both platforms have announced roadmap features for 2024 that will materially affect the comparative evaluation presented in this paper. Engineering teams making platform decisions in early 2024 should monitor the following development trajectories:

- AI-assisted metadata generation: Both OpenMetadata and DataHub have indicated intent to integrate LLM-based description generation that proposes catalog descriptions for undocumented tables and columns - potentially reducing the documentation burden that is the primary barrier to catalog adoption in data engineering teams
- Data contracts standardization: DataHub's Data Contracts feature (beta) and OpenMetadata's emerging SLA and quality commitment primitives are converging toward a common data contract model; the SODA Data Contracts specification [24] and the emerging W3C DCAT standard [25] may provide standardization anchors
- Federated catalog architectures: As data mesh adoption increases, both platforms are investing in federation capabilities that allow decentralized catalog instances to be queried as a unified namespace - directly addressing the organizational challenge of catalog adoption in large enterprises with autonomous business units
- Streaming metadata from SQL engines: BigQuery, Snowflake, and Databricks are exposing richer query-time metadata APIs that enable catalog tools to capture lineage, usage statistics, and schema information in near-real time - potentially closing the lineage freshness gap between pull-based and push-based ingestion models

13.2 Unresolved Research Gaps

This evaluation identifies several dimensions where both platforms exhibit significant limitations that represent active research and development challenges for the metadata management community:

- Semantic lineage: Current lineage models capture structural dependencies (table A feeds table B) but do not capture semantic relationships (column X in table B is the GDPR-relevant personal identifier derived from column Y in table A). Semantic lineage would enable automatic PII propagation and right-to-erasure scope calculation - capabilities that remain largely manual in both platforms
- Cross-cloud catalog federation: Organizations operating data estates across AWS, GCP, and Azure face fragmented catalog instances; neither platform provides a fully satisfactory solution for unified cross-cloud entity resolution with consistent identity management
- Active metadata feedback loops: The vision of catalogs that not only observe but influence pipeline behavior - routing queries to fresher data copies, triggering re-ingestion when staleness is detected, adjusting data product SLAs based on observed quality trends - remains largely aspirational in both platforms' current implementations
- Metadata quality metrics: Paradoxically, neither platform provides robust automated metrics for the quality of the catalog metadata itself - what percentage of tables have descriptions that are accurate (not stale), what fraction of lineage edges are verified against actual data flow rather than inferred from SQL parsing, what is the coverage of ownership assignments across the entity graph

XIV. CONCLUSION

This comparative evaluation has examined OpenMetadata and DataHub across thirteen architectural and operational dimensions, revealing that both platforms have achieved genuine architectural maturity while embodying fundamentally different design philosophies that produce meaningfully different operational experiences.

OpenMetadata's API-first, schema-governed, quality-embedded approach delivers the most complete out-of-box catalog experience for teams that are instrumenting a modern data stack from scratch. Its REST API completeness, native quality execution, and domain-centric data product model align tightly with data mesh organizational patterns and API-native integration workflows. Engineering teams that value catalog consistency, quality observability, and deployment simplicity will find OpenMetadata's design choices well-suited to their requirements.

DataHub's event-driven, graph-native, microservices architecture delivers superior performance at enterprise scale and provides the most mature real-time lineage capture available in the open-source catalog ecosystem. Its Kafka-based metadata event stream, multi-hop graph traversal, and datahub-actions automation framework serve organizations that operate at LinkedIn-scale complexity and require lineage freshness measured in seconds rather than hours. Teams with existing Kafka infrastructure, strong distributed systems expertise, and requirements for 500,000+ entity catalogs will find DataHub's architectural investment in horizontal scalability rewarded.

Critically, the selection between these platforms is not primarily a technology decision but an organizational decision. The platform whose architectural assumptions align with a team's existing infrastructure investments, operational

expertise, and data governance maturity will deliver more value than the platform whose feature list is marginally superior but whose architecture creates friction with the team's working patterns.

As both platforms continue to evolve - with LLM-assisted metadata generation, data contract enforcement, and federated architecture capabilities on the 2024 roadmap - the competitive gap between them will narrow further. Organizations that invest in either platform today will benefit from an improving feature set, an expanding connector ecosystem, and a growing community of practitioners sharing operational knowledge. The open-source data catalog ecosystem has matured to the point where either choice is a credible, enterprise-grade decision - the differentiator is fit, not quality.

KEY FINDING

Final Recommendation: Teams with strong API-integration requirements, embedded quality needs, and data mesh organizational structures should default to OpenMetadata. Teams with large-scale lineage requirements, existing Kafka infrastructure, and event-driven automation needs should default to DataHub. Neither platform is universally superior; organizational context is the decisive factor.

REFERENCES

- [1] OpenLineage Contributors. (2021). OpenLineage: An Open Standard for Metadata and Lineage Collection. Linux Foundation. <https://openlineage.io/spec/>
- [2] Hu, M. et al. (2016). DataHub: A Generalized Metadata Search & Discovery Tool. LinkedIn Engineering Blog. <https://engineering.linkedin.com/blog/2016/08/linkedin-datahub-metadata-search-discovery>
- [3] Collate, Inc. (2021). OpenMetadata: Open Standard for Metadata and Data Collaboration. <https://open-metadata.org/blog/announcing-openmetadata>
- [4] Sawadogo, P. & Darmont, J. (2021). On Data Lake Architectures and Metadata Management. Journal of Intelligent Information Systems, 56(1), 97–120.
- [5] Amundsen Contributors. (2019). Amundsen: Data Discovery and Metadata Engine for Improving the Productivity of Data Analysts, Data Scientists, and Engineers. <https://github.com/amundsen-io/amundsen>
- [6] Apache Atlas PMC. (2020). Apache Atlas: Data Governance and Metadata Framework for Hadoop. Apache Software Foundation. <https://atlas.apache.org/>
- [7] Ehrlinger, L. & Woss, W. (2016). Towards a Definition of Knowledge Graphs. SEMANTiCS (Posters and Demos), 48(1-4), 2.
- [8] Stein, B. & Morrison, A. (2014). The Enterprise Data Catalog: Powering Data Governance in the Digital Age. IBM Institute for Business Value.
- [9] Ghavami, P. (2019). Big Data Governance: Modern Data Management Principles for Hadoop, NoSQL & Big Data Analytics. Amazon Digital Services LLC.
- [10] Rad, B.B., Bhanu, H.J., Ahmadi, M. (2017). An Introduction to Docker and Analysis of its Performance. IJCSNS International Journal of Computer Science and Network Security, 17(3), 228–235.
- [11] Burns, B. et al. (2016). Borg, Omega, and Kubernetes: Lessons Learned from Three Container-Management Systems over a Decade. ACM Queue, 14(1), 70–93.
- [12] Kreps, J., Narkhede, N., Rao, J. (2011). Kafka: A Distributed Messaging System for Log Processing. Proceedings of the NetDB Workshop.
- [13] Taylor, D. et al. (2022). dbt: A New Paradigm for Analytics Engineering. VLDB Workshop on Data Management for End-to-End Machine Learning.
- [14] Kleppmann, M. (2017). Designing Data-Intensive Applications. O'Reilly Media.
- [15] Munir, K. et al. (2022). Metadata Management for the Modern Data Stack. Proceedings of the 2022 International Conference on Information Systems.
- [16] Hai, R., Geisler, S., Quix, C. (2016). Constance: An Intelligent Data Lake System. Proceedings of ACM SIGMOD, 2097–2100.



INTERNATIONAL
STANDARD
SERIAL
NUMBER
INDIA



INTERNATIONAL JOURNAL OF INNOVATIVE RESEARCH

IN COMPUTER & COMMUNICATION ENGINEERING

 9940 572 462  6381 907 438  ijircce@gmail.com



www.ijircce.com

Scan to save the contact details