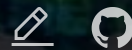


Simple • Sovereign • Scalable Data Stack

Built with DuckDB, dagster and other proven open source components

discover the power of simple. →



Georg Heiler

Solving challenges with data.

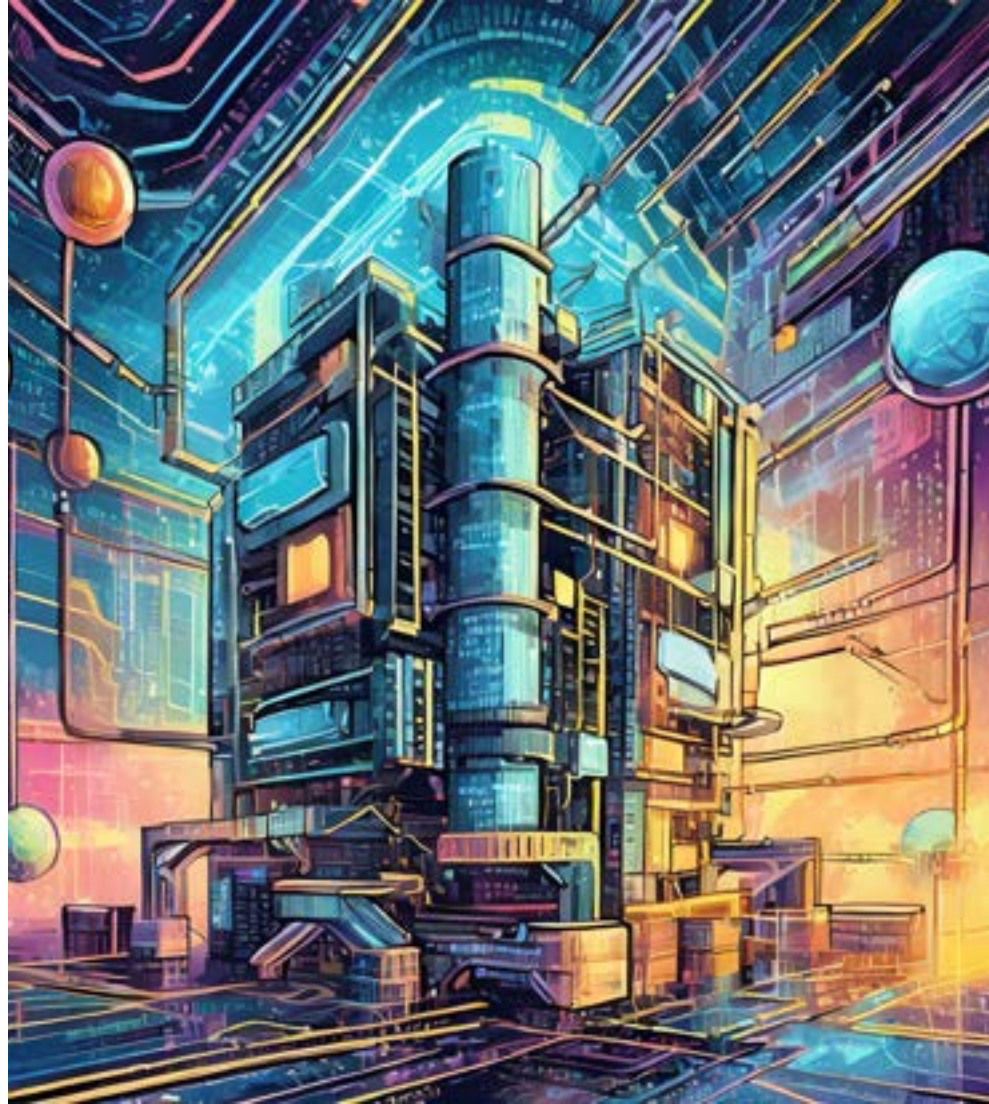
- Senior data expert @Magenta
- Research Software Engineer @ASCII & CSH
- Meetup organizer & frequent Speaker

geoheil.com, linkedin.com/in/geoheil



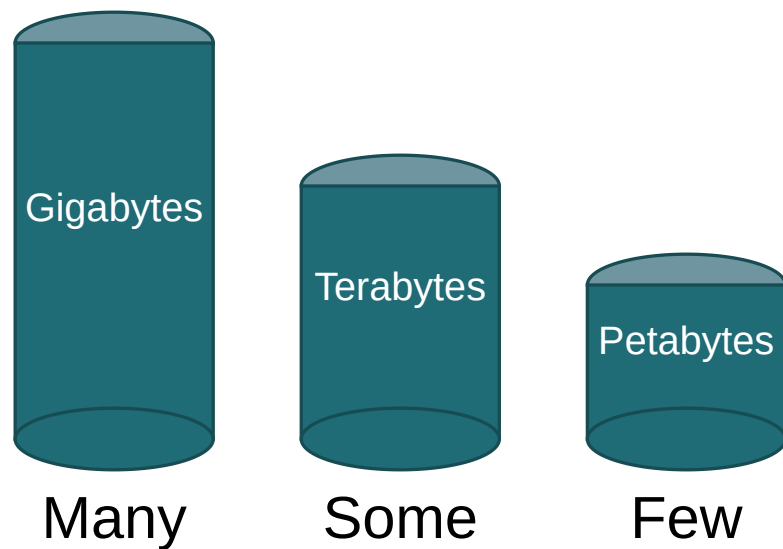
History

- Mainframe
- Data warehouse
- Big Data (Hadoop)
- SQL on large data (Hive, Spark)
- Cloud DWH (Snowflake, BigQuery)
- Cloud exit?
- SQL is there to stay!



Medium-sized data

- Much of the time we work with medium-sized data
- Too big for traditional data processing applications, too small to justify complex data infrastructure



data accumulates over time

Partitions

local data

Reduced latency

Enhanced control (privacy, quality)

Not just laptop; k8s is fine

Where does DuckDB fit?

	Transactional	Analytical
Embedded	SQLite RocksDB	DuckDB
Stand-alone	Oracle Postgres MongoDB	BigQuery Snowflake Starrocks

DuckDB

- Fast, in-process SQL OLAP database
- Vectorized query engine
- Real database (ACID, memory manager)
- Flexible IO: No ingest mandatory (virtualization?)
- Extension ecosystem
 - `spatial`
 - `duckpgq` (graphs)
 - `vectors` (embeddings)
 - `fts` (full text search)
- Innovative and convenient SQL syntax
- Query plan visualizer



Current Cell: [icon]
360,959 Rows
8 Columns [icon]
 [icon] Default [icon]

```
0 Run this query, then click on column names in the right panel to view detailed statistics.
0 */
10
```

type	17	
Sprinter	167k	44%
Thalys	17k	24%
Snelbus	3k	22%
ICE	6k	3%
ICE international	10k	2%
Other	14k	2%
empty seats	71k	1%
Snelbus (p.s. train)	3.78k	10%
Thalys	1.08k	0.2%
ICE international	915	0.2%
Other	1.05k	0.3%

11.190.117	2023-05-15	InterCity	1,410	UT	Ulm	2023-05-15 00:19:00
11.190.117	2023-05-15	InterCity	1,410	GV	Den Haag HS	2023-05-15 00:29:00
11.190.117	2023-05-15	InterCity	1,410	LEDN	Leiden Centraal	2023-05-15 00:45:00
11.190.117	2023-05-15	InterCity	1,410	SHL	Schiphol Airport	2023-05-15 01:03:00
11.190.117	2023-05-15	InterCity	1,410	ASD	Amsterdam Centraal	2023-05-15 01:19:00
11.190.117	2023-05-15	InterCity	1,410	UT	Utrecht Centraal	
11.190.118	2023-05-15	Nightjet	420	NURNB	Nürnberg Hbf	2023-05-15 00:01:00
11.190.118	2023-05-15	Nightjet	420	FFS	Frankfurt (Main) Süd	2023-05-15 01:47:00
11.190.118	2023-05-15	Nightjet	420	FFFFM	Frankfurt (AM) Hbf	
11.190.118	2023-05-15	Nightjet	420	FMF	Frankfurt Flughafen Fernb.	2023-05-15 01:58:00
11.190.118	2023-05-15	Nightjet	420	FMZ	Mann Hbf	2023-05-15 02:16:00
11.190.118	2023-05-15	Nightjet	420	KKG	Koblenz Hbf	2023-05-15 03:13:00
11.190.118	2023-05-15	Nightjet	420	BONN	Bonn Hbf	2023-05-15 04:01:00
11.190.118	2023-05-15	Nightjet	420	KOW	Köln West	

	train_number	station_code	station_name	departure_time	arrival_time
1	573	657			



about 7 days bucketed by day

earliest	2023-05-15	00:13:00.000
latest	2023-05-22	07:14:00.000
low bucket 40	2023-05-22	00:00:00.000
high bucket 54.7%	2023-05-17	00:00:00.000

DuckDB basics

```
1 duckdb
```

DuckDB basics

```
1 CREATE TABLE t1 (  
2   grp VARCHAR,  
3   val INT  
4 );  
5  
6 INSERT INTO t1 (grp, val) VALUES  
7   ('a', 2),  
8   ('a', 1),  
9   ('b', 5),  
10  ('b', 4),  
11  ('a', 3),  
12  ('b', 6);  
13  
14 FROM t1;
```

grp varchar	val int32
a	2
a	1
b	5
b	4

DuckDB basics

```
1  -- find the tok-k largest values per group
2
3  -- traditional window-function based approach
4  SELECT array_agg(rs.val), rs.grp
5  FROM (SELECT val, grp, row_number() OVER (PARTITION BY grp ORDER BY val DESC) AS rid
6        FROM t1
7        ORDER BY val DESC) AS rs
8  WHERE rid < 4
9  GROUP BY rs.grp;
```

10

11

array_agg(rs.val) int32[]	grp varchar
[3, 2, 1]	a
[6, 5, 4]	b

14

15

16

17

array_agg(rs.val) int32[]	grp varchar
[3, 2, 1]	a
[6, 5, 4]	b

DuckDB basics

```
1  -- duckdb: shorter + more efficient
2  SELECT grp, max(val, 3) FROM t1 GROUP BY ALL;
```

grp varchar	max(val, 3) int32[]
a	[3, 2, 1]
b	[6, 5, 4]

```
10
11 -- simpler
12 -- faster
```

DuckDB basics

```
1  -- direct IO, no prior ingestion needed
2  SELECT
3      station_name,
4      count(*) AS num_services
5  FROM 's3://duckdb-blobs/train_services.parquet'
6  GROUP BY ALL
7  ORDER BY num_services DESC
8  LIMIT 10;
```

DuckDB basics

```
1  INSTALL tributary FROM community;
2  LOAD tributary;
3
4  -- Read all messages on a topic
5  SELECT * FROM tributary_scan_topic('test-topic', "bootstrap.servers" := 'kafka-cluster.query.farm:9092');
6
7  -- string similarity (fuzzy matching)
8  INSTALL rapidfuzz FROM community;
9  LOAD rapidfuzz;
10
11 SELECT rapidfuzz_ratio('hello world', 'helo wrld');
```

rapidfuzz_ratio('hello world', 'helo wrld')
double
90.0

DuckDB basics

```
1  ATTACH 'https://github.com/Dtenwolde/duckpgq-docs/raw/refs/heads/main/datasets/snb.duckdb';
2
3  USE snb;
4  install duckpgq FROM community;
5  LOAD
6  duckpgq;
7
8  CREATE OR REPLACE
9  PROPERTY GRAPH snb
10 VERTEX TABLES (
11     Person, Forum
12 )
13 EDGE TABLES (
14     Person_knows_person    SOURCE KEY (Person1Id) REFERENCES Person (id)
15                             DESTINATION KEY (Person2Id) REFERENCES Person (id)
16                             LABEL knows,
17     Forum_hasMember_Person SOURCE KEY (ForumId) REFERENCES Forum (id)
18                             DESTINATION KEY (PersonId) REFERENCES Person (id)
19                             LABEL hasMember
20 );
```

Paradigm shift & limitations

- In-process - No server: Warehouse per user; no mandatory persistent resource allocation
 - No user management
 - Access control via storage layer
 - No dynamic data masking
 - Easily hand a fully hydrated database (with masked data) to 3rd parties
 - Possibility of AWS Lambda/Cloudflare Worker massive scale-out deployments
 - Not only a database - dynamic in-process transformation engine
- Arrow zero-copy integration with (data-science/AI) ecosystem + ADBC for fast BI
 - Arrow Flight (columnar streaming) -> first integration: Boilstream streaming ingestion
 - Columnar layout in memory & SIMD
- Latency reduction - client-side analytics (Rill data); even WASM support
- Single writer multiple readers (locking)

pg_duckdb

- Remember: powerful in-process transformation engine but limitation of single writer
- Combine the strengths of postgres (proven, IAM) as a serving layer with the power of duckdb

```
-- remember DuckDB solo
-- CREATE TABLE stations AS FROM 's3://duckdb-blobs/s

-- combination
CREATE TABLE pg_stations AS
select *
from duckdb.query('SELECT * FROM ' 's3://duckdb-blobs/

-- postgres
select *
from pg_stations
limit 10;
```



System Catalog

OLAP
Compute

Storage Engine

Table Format

File Format (column-oriented)

Storage

Catalog

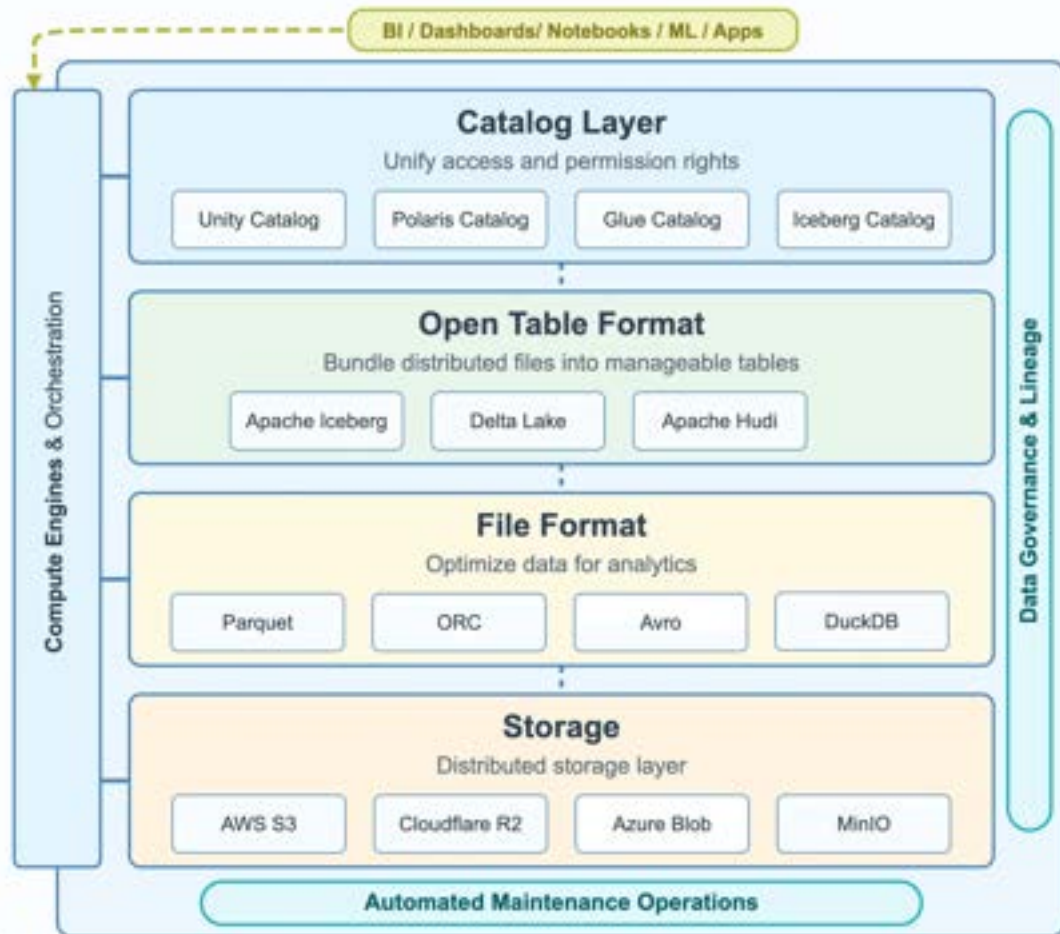
Compute

File Format

Storage

Open Data Platform Architecture

Built on Open Standards and Formats



Iceberg catalog

- S3 until recent - no strong consistency
- Read-after-write consistency was missing
- Impact: Add layer to handle transactions

TLDR: Add a database in disguise

Less features + custom API; no SQL.

- No global transaction.
- No neat SQL filters.
- Lots of complexity

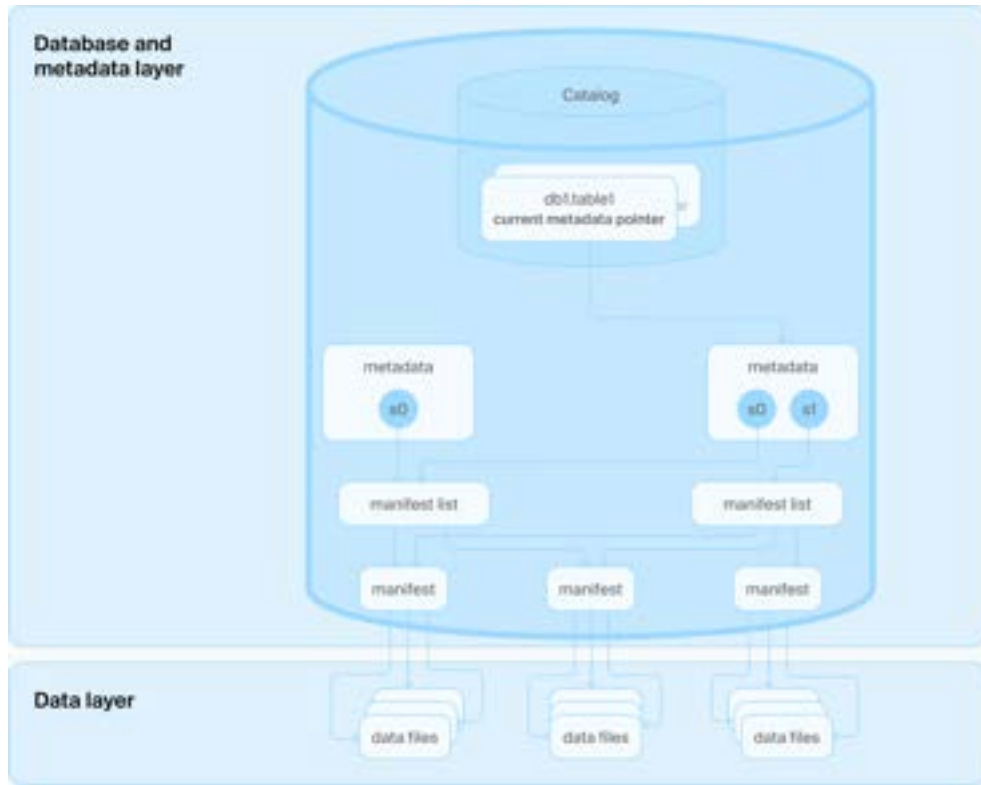
Problem: Table format was specifically designed to not require a database



Ducklake

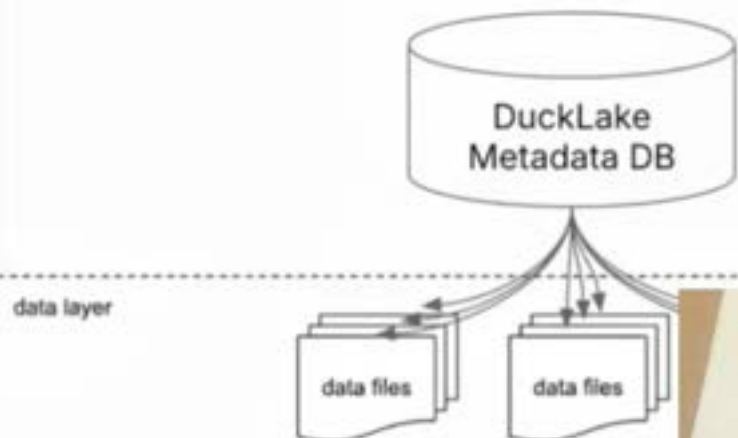
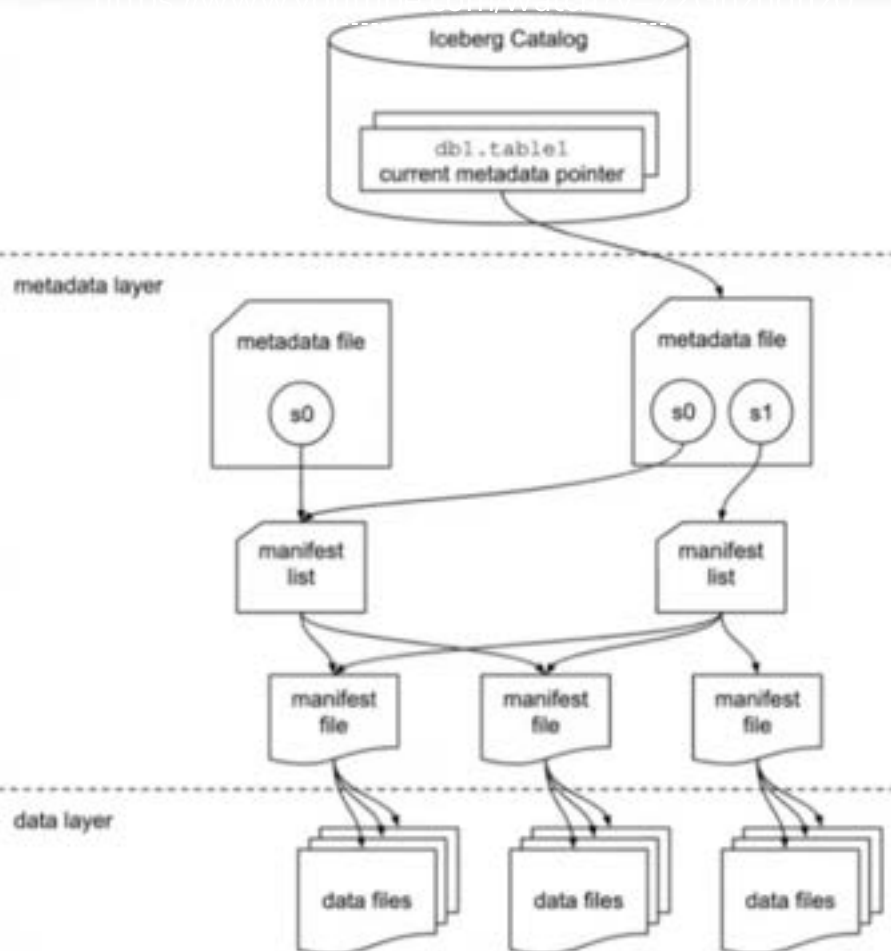
- Database and SQL first design
- Open standard
- Flexible backends
 - Catalog (postgres, sqlite, duckdb, mysql)
 - Storage (local fs, s3 compatible, ...)
- Real ACID transactions (not just for a single table)
- Reduce complexity (increase performance - inlining)
- Bridge gap to table formats
- Simple encryption of data

Most cloud data warehouses internally have a similar architecture!





Replacing a whole lot of stuff with a database



Ducklake basics

```
1 duckdb
```

Ducklake basics

```
1  INSTALL ducklake;
```

Ducklake basics

```
1  ATTACH 'ducklake:metadata.ducklake' AS my_ducklake;  
2  CREATE TABLE my_ducklake.demo  
3  (  
4      i INTEGER  
5  );  
6  INSERT INTO my_ducklake.demo  
7  VALUES (42),  
8          (43);  
9  FROM my_ducklake.demo;
```

10

11

12

13

14

15

16

17

i
int32
42
43

Ducklake basics

```
1 DELETE
2 FROM my_ducklake.demo
3 WHERE i = 43;
4 FROM my_ducklake.demo;
```

5

6

i
int32
42

10

11

Ducklake basics

```
1 BEGIN TRANSACTION;  
2 DELETE  
3 FROM my_ducklake.demo;  
4  
5 FROM my_ducklake.demo;
```

i
int32
0 rows

```
12  
13 ROLLBACK;  
14 FROM my_ducklake.demo;
```

i
int32
42

Ducklake basics

```
1 FROM ducklake_snapshots('my_ducklake');
```

snapshot_id int64	snapshot_time timestamp with time zone	schema_version int64	changes map(varchar, varchar[])
0	2025-08-31 13:01:19.451+02	0	{schemas_created=[main]}
1	2025-08-31 13:01:24.263+02	1	{tables_created=[main.demo]}
2	2025-08-31 13:01:24.511+02	1	{tables_inserted_into=[1]}
3	2025-08-31 13:01:33.954+02	1	{tables_deleted_from=[1]}

```
10
```

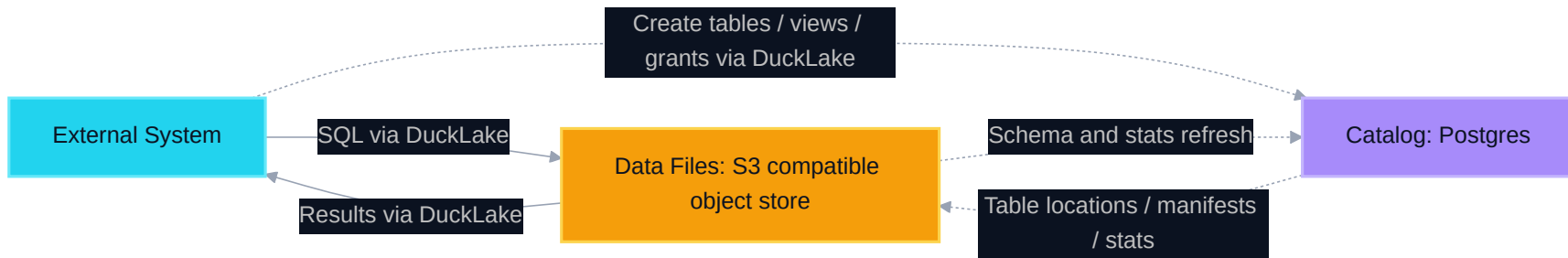
```
11 FROM my_ducklake.demo AT (VERSION ⇒ 2);
```

i int32
42
43

```
18
```

Ducklake: Example architecture

- Catalog: Postgres
- Storage: S3 compatible (MinIO) for parquet files
- All locally in docker with SSL + DNS



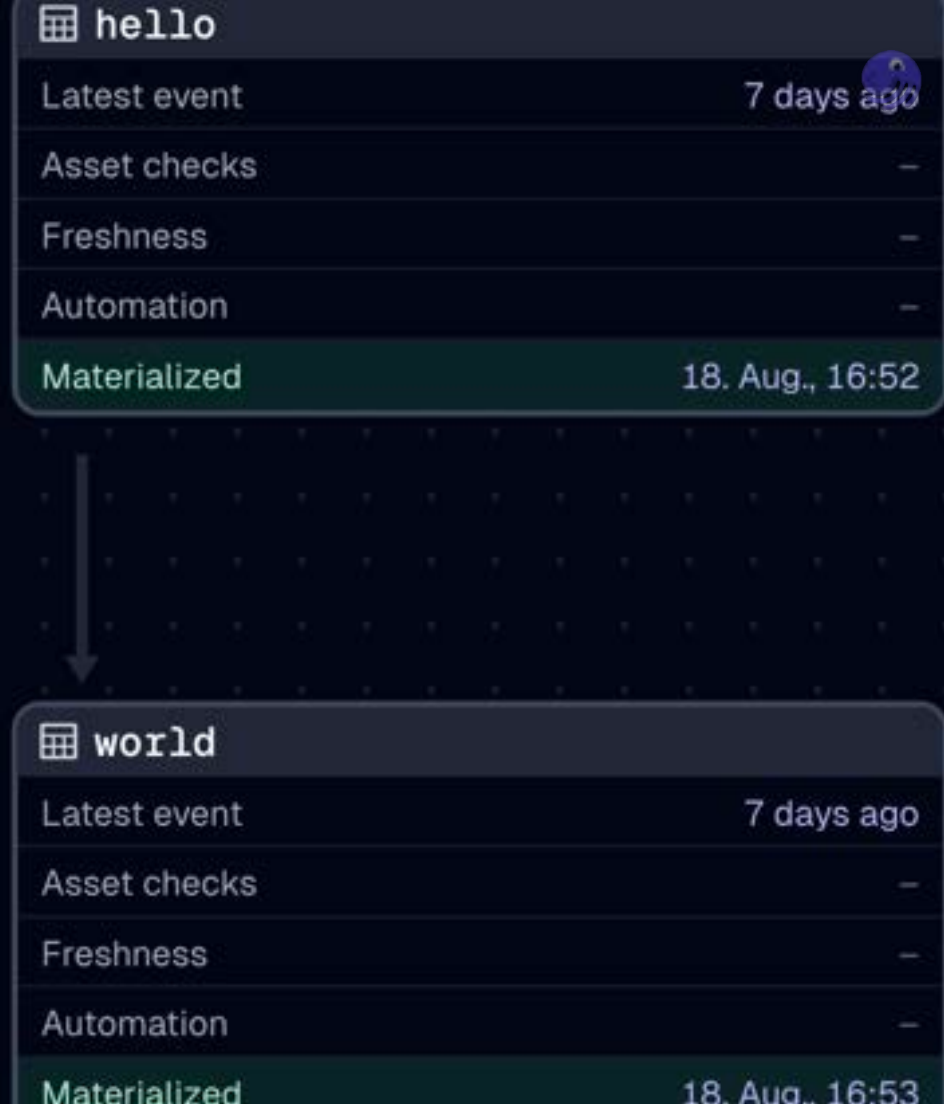
Dagster & asset graph

- Like a calculator for crunching numbers
- Graph allows computer to reason about data dependencies
- Rapid iteration: Just edit code. No need to wait for XYZ SaaS service
- Break down tool and department silos

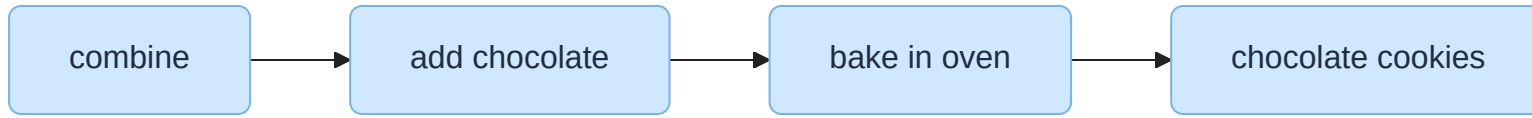
```
import dagster as dg

@dg.asset
def hello(context: dg.AssetExecutionContext):
    context.log.info("Hello!")

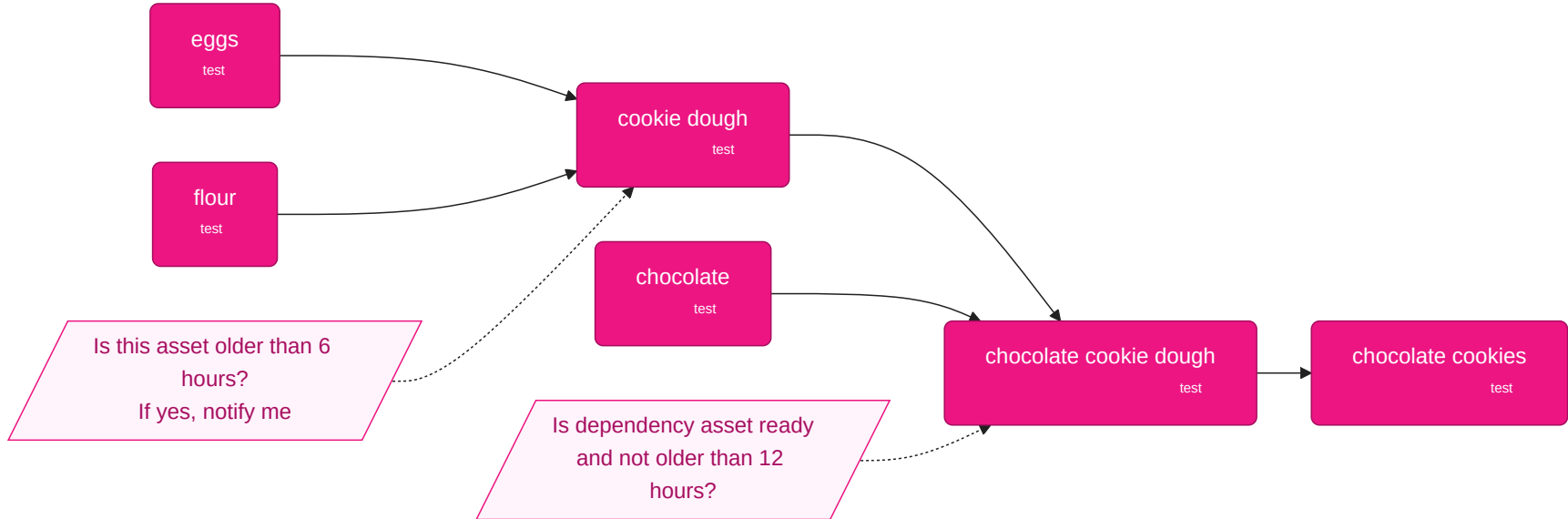
@dg.asset(deps=[hello])
def world(context: dg.AssetExecutionContext):
    context.log.info("World!")
```



Task-based orchestrator



Asset-based orchestration



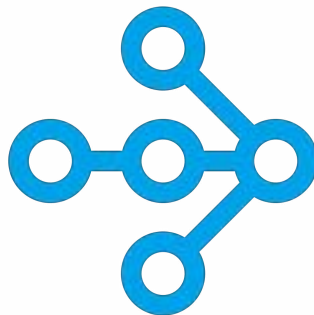
Process automation: Small Language Models

- No Human to chat and correct
- Cost and energy efficient
- Minimal latency
- Many small AI enabled steps with context engineering for repeatable reliable results



Ray

- Easily scaling arbitrary computation graphs
- Minimal devops support needed
- Seamless scaling: From laptop to cluster, no code rewrite.
- Elastic & cost-efficient: Scale up/down with demand, minimize idle GPUs.
- Batch and API serving support



RAY

Combined showcase

Ducklake (postgres catalog, MinIO as object storage)

Case 1: `dlt` ducklake

- REST API
- Ingestion to ducklake
- Dagster integration

Case 2: AI - Small language models

For ASR and document processing via docling

- V1: Plain python
- V2: Scaled with ray
- V3: Ray integration with dagster



Demo 1: dlt ducklake ingestion

```
1 # OSS dagster ducklake integration to follow along
2 # https://github.com/dagster-io/community-integrations/tree/main/libraries/dagster-ducklake
3
4 import dagster as dg
5
6 @dg.asset
7 def my_ducklake_asset(ddb: DuckDBConnectionProvider):
8     with ddb.duckdb_connect() as con:
9         query = "select * from table"
10
11         df = con.query(query).df()
12         df.head()
```

Demo 1: dlt ducklake ingestion

```
1  # define resources
2  DuckLakeResource(
3      metadata_backend=PostgresConfig(
4          host="db.duckpond.data-engineering.expert",
5          port=5433,
6          database="ducklake_catalog",
7          user=dg.EnvVar("POSTGRES_USER"),
8          password=dg.EnvVar("POSTGRES_PASSWORD"),
9      ),
10     storage_backend=S3Config(
11         endpoint_url="objectstore.duckpond.data-engineering.expert",
12         bucket="duckpond-dev",
13         prefix="stage",
14         aws_access_key_id=dg.EnvVar("OBJECT_STORE_ROOT_USER"),
15         aws_secret_access_key=dg.EnvVar("OBJECT_STORE_ROOT_PASSWORD"),
16         region=dg.EnvVar("OBJECT_STORE_REGION"),
17         use_ssl=True,
18         url_style="path",
19     ),
20     alias="stage",
21     plugins=["postgres", "httpfs", "ducklake"],
22 )
```

Demo 1: dlt ducklake ingestion

```
1 # prepare sqlalchemy connection
2
3 @contextmanager
4 def duckdb_connect(self) → Generator[DuckDBPyConnection, None, None]:
5     """Yields a pre-configured native DuckDB connection for safe use."""
6     conn = self.get_duckdb_connection()
7     # ...
8     # CREATE OR REPLACE SECRET {ducklake_secret_name} (TYPE DUCKLAKE ...
9     # ATTACH 'ducklake
10    # ...
11    try:
12        yield conn
13    finally:
14        conn.close()
```

Demo 1: dlt ducklake ingestion

```
1 # use with dlt
2 pokemon_source = rest_api_source({
3     "client": {"base_url": "https://pokeapi.co/api/v2/"},
4     ...
5
6 POKE_RESOURCES = {"pokemon": "General Pokemon data, loaded to Duck Lake.",}
7
8 poke_asset_specs = [
9     dg.AssetSpec(
10         key=["poke_api_ducklake", name],
11         description=desc,
12         skipable=True,
13         kinds={"duckdb", "dlt"},
14     )
15     for name, desc in POKE_RESOURCES.items()
16 ]
```

Demo 1: dlt ducklake ingestion

```
1  @dgt.multi_asset(  
2      specs=poke_asset_specs,  
3      group_name="pokemons",  
4      can_subset=True,  
5  )  
6  def load_pokemon_to_ducklake(  
7      context: dg.AssetExecutionContext, ducklake: DuckLakeResource  
8  ):  
9      selected_resource_names = {key.path[-1] for key in context.selected_asset_keys}  
10     source_to_run = pokemon_source.with_resources(*selected_resource_names)  
11     pipeline = dlt.pipeline(  
12         pipeline_name="rest_api_pokemon_to_ducklake",  
13         dataset_name="poke_rest_api",  
14         destination=dlt.destinations.sqlalchemy(ducklake.get_engine()),  
15     )  
16     load_info = pipeline.run(source_to_run)
```

Demo 2: AI small language models + ray (ASR)

```
1 # huggingface whisper large 3 turbo
2 # https://huggingface.co/openai/whisper-large-v3-turbo
3
4 class UniversalASRTranscriber:
5     def __init__(self, model_id: str):
6         self.model_id = model_id
7         self._init_whisper()
8     def _init_whisper(self):
9         self.model = AutoModelForSpeechSeq2Seq.from_pretrained(
10             self.model_id, dtype=self.dtype, low_cpu_mem_usage=True
11         )
12         self.model.to(self.device)
13         self.processor = AutoProcessor.from_pretrained(self.model_id)
14         self.pipe = pipeline(
15             "automatic-speech-recognition",
16             model=self.model,
17             ...
18         )
```

Demo 2: AI small language models + ray (ASR)

```
1  def transcribe_one(  
2      self, audio_path: str,  
3  ):  
4      audio, sr = load_audio_16k_mono(audio_path, target_sr=16000)  
5      out = self.pipe(audio)  
6      text = out.get("text", str(out)) if isinstance(out, dict) else str(out)  
7      return {  
8          "path": audio_path,  
9          "text": text,  
10     }
```

Demo 2: AI small language models + ray (ASR)

```
1 model_id = "openai/whisper-large-v3-turbo"
2 transcriber = UniversalASRTranscriber(model_id=model_id)
3
4 transcriber.transcribe_one("path/to/audio/file.mp3")
```

Demo 2: AI small language models + ray (ASR)

```
1 class RayASRWorker:
2     def __init__(self, model_id: str):
3         self.transcriber = UniversalASRTranscriber(model_id=model_id)
4
5     def transcribe(self, audio_path: str, **kwargs) → Dict[str, Any]:
6         return self.transcriber.transcribe_one(audio_path, **kwargs)
7
8 def build_asr_mapper(*, model_id: str):
9     class _ASRMapper:
10         def __init__(self):
11             self.worker = RayASRWorker(model_id=model_id)
12
13         def __call__(self, batch):
14             results = []
15             for audio_path in batch["path"]:
16
17                 result = self.worker.transcribe(audio_path, **kwargs)
18                 results.append(result)
19
20             return {"result": results}
21
22     return _ASRMapper
```

Demo 2: AI small language models + ray (ASR)

```
1 def run_asr_on_ray_cluster(  
2     model_id: str,  
3     inputs: List[str],  
4     num_workers: int,  
5     logger,  
6 ):  
7     ds = rd.from_items([{"path": p} for p in inputs])  
8  
9     result_ds = ds.map_batches(  
10         build_asr_mapper(model_id=model_id),  
11         batch_size=4,  
12         num_cpus=1, num_gpus=1, resources={"MY_AI_NODE": 1},  
13     )  
14  
15     materialized = result_ds.materialize()
```

Demo 2: AI small language models + ray (ASR)

```
1 ray.init(resources={"MY_AI_NODE": 1})
2
3 all_results = []
4 for model_id in ["openai/whisper-large-v3-turbo"]:
5     results = run_asr_on_ray_cluster(
6         model_id=model_id,
7         inputs=["data/obama.mp3",],
8     )
9     all_results.extend(results)
```

Demo 2: AI small language models + ray (ASR)

```
1  @dg.asset
2  def asr_ray_integrated(
3      context: dg.AssetExecutionContext,
4      ray_cluster: RayResource,
5  ):
6      inputs = ["/ai_example/data/obama.mp3"]
7      model_id = "openai/whisper-large-v3-turbo"
8
9      results = run_asr_on_ray_cluster(
10         model_id=model_id,
11         inputs=inputs,
12     )
13
14     total_duration = sum(r["duration_seconds"] for r in results)
15     context.add_output_metadata(
16         {
17             "num_files_processed": len(results),
18             "total_transcription_time": format_duration(total_duration),
19             "model_id": model_id,
20             "first_result": results[0]["text"] if results else "N/A",
21         }
22     )
23     return results
```

Demo 3: Sovereign document intelligence

```
1 from docling.datamodel.base_models import InputFormat
2 from docling.document_converter import DocumentConverter
3
4 source = "data/W02021041671-eval-small.pdf"
5
6 class MyDocumentConverter:
7     def __init__(self, document_format: InputFormat = InputFormat.PDF):
8         self.converter = DocumentConverter()
9         self.converter.initialize_pipeline(document_format)
10
11     def convert_one(self, uri: str):
12         return self.converter.convert(uri).document
13
14 converter = MyDocumentConverter()
15 doc = converter.convert_one(source)
16
17 print(doc.export_to_markdown())
```

Demo 3: Sovereign document intelligence

```
1 class DoclingActor:
2     def __init__(self):
3         self.converter = MyDocumentConverter()
4
5     def convert_one(
6         self, path: str, out_dir: str,
7     ) → ConvertSummary:
8         res: ConversionResult = self.converter.convert_one(src_path)
9
10        # handle persistence from distributed actor
11        handler = markdown_referenced_only_handler
12        info = handler(
13            src_path=src_path,
14            result=res,
15            out_dir=Path(out_dir),
16        )
```

Demo 3: Sovereign document intelligence

```
1 def build_converter(actor_kwargs: Dict[str, Any], out_dir: str):
2     class _DocConverter:
3         def __init__(self):
4             self.worker = DoclingActor(**actor_kwargs)
5             self.out_dir = out_dir
6
7         def __call__(self, batch: Dict[str, List[str]]) → Dict[str, List[Any]]:
8             rows: List[Dict[str, Any]] = []
9             for p in batch["path"]:
10                 s = self.worker.convert_one(
11                     src_path=p,
12                     out_dir=self.out_dir,
13                 )
14                 rows.append( ... result statistics ... )
15
16             cols = {k: [r[k] for r in rows] for k in rows[0].keys()}
17             return cols
18     return _DocConverter
```

Demo 3: Sovereign document intelligence

```
1 files = [str(p) for p in Path().glob(INPUT_GLOB)]
2 ray.init(resources={"DOCL_NODE": 1})
3
4 ds = rd.from_items([{"path": p} for p in files])
5
6 result = ds.map_batches(
7     build_converter(out_dir=OUTPUT_DIR),
8     num_cpus=4,
9     num_gpus=1,
10    resources={"DOCL_NODE": 1},
11 )
12
13 status_ds = result.materialize()
```

Demo 3: Sovereign document intelligence

```
1  @dg.asset
2  def document_ray_integrated(
3      context: dg.AssetExecutionContext,
4      ray_cluster: RayResource,
5  ) → Dict[str, Any]:
6      data_dir = "/path/to/data"
7      pattern = "*.pdf"
8
9      files = [str(p) for p in data_dir.glob(pattern)]
10     ds = rd.from_items([{"path": p} for p in files])
11
12     mapper_document = build_converter(out_dir="/path/to/output")
13     result = ds.map_batches(
14         mapper_document,
15         num_cpus=4,
16         num_gpus=1,
17         resources={"DOCL_NODE": 1},
18     )
19     status_ds = result.materialize()
20
21     total = status_ds.count()
22
23     return { ... statistics of results, processed records, failures ... }
```

Ray resource configuration

```
from dagger_ray import LocalRay
from dagger_ray.kuberay import in_k8s, KubeRayInteractiveJob

ray_cluster = LocalRay() if not in_k8s else KubeRayInteractiveJob()

definitions = dg.Definitions(
    assets=[my_distributed_computation],
    resources={"ray_cluster": ray_cluster},
)
```

Full ray k8s JSON can be specified as needed.

Futher considerations

Do not return distributed IO to orchestrator

Consider a mutli-modal feature graph engine

Daniel & Georg are developing one

- Row based lineage tracking
- Recompute only what is needed
- Rapid experimentation through memoization

Reach out if you want to learn more (stil early WIP)

Takeaway

- Ducklake is simple, powerful and flexible
- Ecosystem maturity (integrations) needs more time
- Asset based orchestration solves challenges:
 - Inside of datavault for more speed
 - Outside/around DV for AI/ML and ingestion
 - Not stuck on a single execution engine
 - Encapsulate complexity: Easy vibe-coding or onboarding of less technical users
 - Advanced: Environment-re-targeting possible to save money

Local-first data ecosystem is possible, and has advantages

Closing summary

Take back control: Reduce entropy

Systems should scale as needed:

- Down to a single developer`s notebook
- As big as necessary, but with simplicity

You should be in control!

geoheil.com

Dependency handling

- Conda support is a must
- Conda itself is slow and dated (even with mamba)

`pixi` as the solution

- Fast
- Lockfiles
- Multi-environment handling (dev, prod)
- Pip and conda support
- Neatly packaging of environments
- Easy bootstrap
- Task runner



Security

- Sops (multiple backends supported: age, various key stores)
- Age (pgp like without the mess)

Simple: `.env` file with secrets

- Encrypt
- Push to git - or your kms store of choice

Ducklake demo

```
1  -- setup ducklake
2  INSTALL ducklake;
3  INSTALL postgres;
4  INSTALL httpfs;
5  LOAD httpfs;
```

Ducklake demo

```
1  -- define secrets
2  CREATE OR REPLACE SECRET duckpond_secret_catalog (
3      TYPE postgres,
4      HOST 'db.duckpond.data-engineering.expert',
5      PORT 5433,
6      DATABASE ducklake_catalog,
7      USER getenv('CATALOG_DB_USER'),
8      PASSWORD getenv('CATALOG_DB_PASSWORD')
9  );
10
11 CREATE OR REPLACE SECRET duckpond_secret_storage (
12     TYPE S3,
13     KEY_ID      getenv('OBJECT_STORE_ROOT_USER'),
14     SECRET      getenv('OBJECT_STORE_ROOT_PASSWORD'),
15     ENDPOINT    'objectstore.duckpond.data-engineering.expert',
16     URL_STYLE   'path',
17     REGION      getenv('OBJECT_STORE_REGION'),
18     USE_SSL     true,
19     SCOPE       's3://my-duck-lake'
20 );
```

Ducklake demo

```
1  -- conect
2  CREATE OR REPLACE SECRET duckpond_ducklake (
3      TYPE DUCKLAKE,
4      METADATA_PATH '',
5      DATA_PATH 's3://my-duck-lake/',
6      METADATA_PARAMETERS MAP {'TYPE': 'postgres', 'SECRET': 'duckpond_secret_catalog'}
7  );
8  ATTACH 'ducklake:duckpond_ducklake' AS duckpond;
9
10 USE duckpond;
```

Ducklake demo

```
1 CREATE TABLE nl_train_stations AS FROM 'https://blobs.duckdb.org/nl_stations.csv';
```

```
2
```

```
3 SELECT * FROM (SHOW ALL TABLES) WHERE schema = 'main';
```

```
4
```

database varchar	schema varchar	name varchar	column_names varchar[]	column_types varchar[]
duckpond	main	nl_train_stations	[id, code, uic, na...	[BIGINT, VARCHAR, BIGINT, VARCHAR, VARCHAR, VARCHAR, VARCHAR...

```
9
```

Ducklake demo

```
1 mc ls duckpond/my-duck-lake/main/nl_train_stations
2 # 1 file
3
4 [2025-09-29 19:53:52 CEST] 41KiB STANDARD ducklake-0199969c-0a24-7b90-a7fd-562c7a2e4d69.parquet
```

Ducklake demo

```
1 UPDATE nl_train_stations SET name_long='Johan Cruijff ArenA' WHERE code = 'ASB';  
2 SELECT name_long FROM nl_train_stations WHERE code = 'ASB';
```

3

4

5 name_long

6 varchar

7

8 Johan Cruijff ArenA

9

Ducklake demo

```
1 mc ls duckpond/my-duck-lake/main/nl_train_stations
2 # More files current, one delta, one deletion vector
3
4 [2025-09-29 19:53:52 CEST] 41KiB STANDARD ducklake-0199969c-0a24-7b90-a7fd-562c7a2e4d69.parquet
5 [2025-09-29 19:56:58 CEST] 1.8KiB STANDARD ducklake-0199969e-df9a-74ab-9f40-9542612bbc0e.parquet
6 [2025-09-29 19:56:58 CEST] 790B STANDARD ducklake-0199969e-dfb7-7267-b817-ff61d6edaba4-delete.parquet
```

Relative comparison

of analytical engines

- ELO score: Relative not absolute
- Robustness: Continuously improve ranking, robust to cherry picking

See more at: <https://georgheiler.com/post/elo-data-challenge/>

1	 *StarRocks	2017		
2	 *Kinetica	2015		
3	 *Apache Hudi	2012	-2	
4	 *Databricks SQL	2005	-5	
5	 *DuckDB	2004	-20	
6	 *Clickhouse	2004	+19	 
7	 *Delta Lake	1985	-2	
8	 *TrinoDB	1985		
9	 *Snowflake	1981	-5	
10	 *Apache Spark	1975		

Useful OSS projects

- local data stack template
- dagster-vertexai integration
- dagster-ducklake integration
- dagster-tableau integration