



Adrian Bishop

Jr-Mid Level GCP Data Engineer

📞 (602) 555-0147 ✉️ adrian.bishop@example.com

🌐 linkedin.com/example/adrianbishop 📍 Phoenix, AZ 85016

SUMMARY

Mid-level **GCP Data Engineer** with more than five years building batch and real-time data platforms across Hadoop and Google Cloud environments. Delivers scalable PySpark and Python pipelines, Kafka streaming flows, and governed data lake migrations that support analytics and operational teams. Translates business requirements into technical designs, integration tests, deployment plans, and reliable production workloads. Hands-on work spans BigQuery, **Cloud Storage**, Dataproc, **Cloud Composer**, **Spark Streaming**, Flink, schema validation, access management, and change management. Career growth reflects increasing ownership, from supervised Hadoop assignments through platform migration leadership and architecture partnership. Collaborative delivery with engineers, architects, and stakeholders keeps modernization practical while protecting existing operations. Ready to help a Phoenix team move critical workloads into GCP with dependable pipelines, disciplined releases, and clear platform support.

EXPERIENCE

Senior Data Engineer

Mesa Verde Analytics 📅 July 2024 - Present 📍 Phoenix, AZ

Owns migration-ready data pipelines for retail operations, supporting Hadoop continuity and incremental GCP adoption. Partners with architects and analytics stakeholders across design, testing, deployment, access review, and production support.

- Built **Dataproc** PySpark pipelines moving inventory data into **BigQuery** with source reconciliation controls.
- Created Kafka and Spark Streaming flows with checkpointing and replay for order event recovery.
- Orchestrated Cloud Composer dependencies through **Cloud Storage landing zones** and service account controls.
- Led **integration testing**, release checklists, access reviews, and production change-management preparation.
- Tuned Spark resources, partitioning, and **file formats** across **Hadoop** and GCP workloads for steadier processing.
- Coordinated migration designs with platform architects and analytics stakeholders for dependable retail data delivery.

Data Engineer

Copper State Data Works 📅 March 2022 - June 2024 📍 Tempe, AZ

Built Hadoop and early GCP data services for customer, transaction, and product domains. Converted requirements into governed pipeline designs while partnering with architects and infrastructure teams on releases, reliability, and access controls.

- Built Spark pipelines on Hadoop for customer, transaction, and product data supporting six analytics groups.
- Converted requirements into Python and **PySpark** designs for governed **data lake** ingestion and publication.
- Implemented Kafka consumers and Spark Streaming jobs with **dead-letter handling** and **operational logging**.
- Loaded approved Hadoop extracts into Cloud Storage and BigQuery during early GCP **platform enablement**.
- Created Cloud Composer workflows with integration tests and **deployment runbooks** for repeatable releases.
- Resolved failed jobs, large joins, and permissions issues with architects and infrastructure teams.

Associate Data Engineer

Sonoran Information Services 📅 November 2020 - February 2022 📍 Scottsdale, AZ

STRENGTHS



Migration planning

Mapped Hadoop dependencies before each move. That preparation gave architects clearer choices and helped reporting teams trust staged GCP releases.



Streaming reliability

Added checkpointing, replay, and validation during order-event work. Teammates gained a calmer recovery path when downstream systems interrupted processing.



Pipeline design

Turned business needs into practical Python and PySpark flows. Stakeholders saw requirements become testable delivery plans with visible data controls.



Team partnership

Worked closely with architects, infrastructure teams, and analysts. Colleagues often brought difficult joins and permissions issues for collaborative problem solving.



Release discipline

Used integration tests, access reviews, and change records before production releases. That rhythm made deployments easier for platform teams to support.

SKILLS

[Python](#) [PySpark](#) [Hadoop](#)

[Spark Streaming](#) [Kafka](#) [Flink](#)

[BigQuery](#) [Cloud Storage](#) [Dataproc](#)

Cloud Composer Data lakes

Batch processing

Real-time processing

Data migration

Pipeline orchestration

LANGUAGES

English Native

Spanish Intermediate

MY CAREER



● Senior Data Engineer at Mesa Verde Analytics (2.2 Years)

● Data Engineer at Copper State Data Works (2.2 Years)

● Associate Data Engineer at Sonoran Information Services (1.2 Years)

Supported supervised Hadoop engineering assignments and enterprise reporting feeds across finance and customer operations. Developed foundational Python, PySpark, Kafka, testing, monitoring, and migration assessment skills while presenting findings to senior engineers and analysts.

- Wrote **Python** and PySpark jobs applying quality rules and publishing datasets with documented schemas.
- Configured **Kafka** ingestion patterns and tested Spark Streaming latency, restart behavior, and **duplicate handling**.
- Profiled on-premise datasets and cataloged dependencies for future **Cloud Storage** and BigQuery migrations.
- Executed unit and integration tests with defect evidence and deployment documentation for reviewed releases.
- Monitored Hadoop jobs, analyzed **failed stages**, and presented remediation findings to senior engineers.
- Supported **reporting feeds** across finance and customer operations through supervised pipeline assignments.

PROJECTS

Retail Data Platform Migration 📅 2025

Partnered with platform architects and analytics users on a staged Hadoop migration into GCP. Built reconciliation checks, clarified ownership, and kept reporting teams informed through each release.

Streaming Reliability Initiative 📅 2024

Worked with engineers on Kafka and Spark Streaming recovery patterns. Shared replay procedures and operational findings, helping teammates restore event flows with less disruption.

LEADERSHIP & AWARDS

- Dean's List, Arizona State University, 2019
- Outstanding Data Systems Capstone Award, Arizona State University, 2020
- Hadoop and Cloud Migration Project Recognition, Sonoran Information Services, 2021

EDUCATION

Bachelor's Degree in Computer Information Systems

Arizona State University 🎓 GPA: 3.8 📅 2020 📍 Tempe, AZ

Coursework: Data Engineering, Database Systems, Cloud Computing, Systems Analysis

CERTIFICATIONS

- Google Cloud Professional Data Engineer 📅 2025
- Google Cloud Digital Leader 📅 2023

TECHNICAL SKILLS

- **Programming Languages:** Python, SQL, PySpark
- **Hadoop Ecosystem:** Hadoop, HDFS, YARN
- **Google Cloud Data:** BigQuery, Cloud Storage, Dataproc
- **Workflow Orchestration:** Cloud Composer, Apache Airflow, DAGs
- **Stream Processing:** Kafka, Spark Streaming, Flink
- **Data Processing:** Apache Spark, Spark SQL, PySpark
- **Data Engineering:** ETL, ELT, batch processing
- **Data Lake Systems:** Data lakes, landing zones, curated zones
- **Testing Practices:** Unit testing, integration testing, schema validation
- **Deployment Practices:** Production deployment, release checklists, runbooks
- **Governance Controls:** Change management, access management, service accounts
- **Data Quality:** Data quality, reconciliation, duplicate handling

PROFESSIONAL AFFILIATIONS

- Data Engineering Peer Mentor, Arizona State University, 2019 - 2020
- Cloud Computing Study Group Coordinator, Arizona State University, 2019 - 2020
- Volunteer Technical Presenter, Phoenix Data Professionals Meetup, 2023 - Present

ADDITIONAL INFORMATION

Work Status : Authorized to work in United States. No sponsorship required.

REFERENCES

AVAILABLE ON REQUEST