

ABDUL BASEER SHAIK

DATA ENGINEER

+1 240-245-0114 | abdulbaseershaiks@gmail.com | [linkedin.com/in/abdulbaseershaik](https://www.linkedin.com/in/abdulbaseershaik)

PROFESSIONAL SUMMARY

- Data Engineer with **12+** years migrating legacy **Oracle** and **Hadoop** workloads onto cloud-native stacks without breaking reconciliation, across finance, government, and healthcare where a failed control becomes an audit finding.
- Led migration of **45+** production ETL workloads to **AWS Glue**, EMR, **Redshift**, and **Snowflake** across financial services, government, healthcare, and retail.
- Removed **4+** hours from nightly batch windows and **\$310K** from annual infrastructure spend through **Spark** tuning, partitioning strategy, and cluster right-sizing.
- Certified across all four platforms (Fabric **DP-600**, **Power BI PL-300**, **Azure DP-203**, **Databricks Data Engineer Professional**, **Snowflake SnowPro**, AWS AI Practitioner), administering Fabric tenant configuration, workspace architecture, **F-SKU** capacity management, and environment separation with **CI/CD** deployment pipelines.
- Layered Bronze, Silver, and Gold medallion zones with schema validation, cleansing, normalization, and reconciliation testing at every boundary.
- Delivered across AWS (Glue, S3, EMR, **Redshift**, Lambda, Step Functions, Kinesis), **Azure** (Data Factory, **Databricks**, Synapse, **ADLS Gen2**), and **GCP** (**BigQuery**, **Dataproc**, **Dataflow**, Composer).
- Specialized in **PySpark** and **Spark SQL** performance tuning: partitioning, broadcast joins, caching strategy, executor sizing, and skew mitigation.
- Tuned **Snowflake** and **Amazon Redshift** through warehouse sizing, clustering keys, distribution and sort key selection, workload management queues, and Spectrum external tables.
- Engineered **Databricks** workloads on **Delta Lake** and **Delta Live Tables**, enforcing row-level quality with **DLT** expectations and governing access through **Unity Catalog** across regulated datasets.
- Instituted data-quality and reconciliation controls that withstood audit, covering source-to-target completeness, control totals, exception routing, and restart logic.
- Orchestrated multi-stage pipelines in **Amazon MWAA** and **Apache Airflow** with dependency management, retries, alerting, and backfill support.
- Shipped batch and near-real-time ingestion with **Amazon Kinesis**, **Kafka**, and **Spark Streaming**, including topic partitioning, consumer group scaling, and offset management.
- Built LLM-assisted data workflows with **LangGraph**, structured prompt engineering, agent tool-calling, bounded retry loops, and schema-validation guardrails.
- Modeled Star and **Snowflake** schemas, Slowly Changing Dimensions Type 1 and 2, Data Marts, ODS, and enterprise data warehouse layers.
- Processed Parquet, ORC, Avro, JSON, CSV, and Sequence formats with format-appropriate compression and partitioning strategy.
- Optimized **Hive** through partitioning, bucketing, join strategy, UDF development, and **HiveQL** migration from legacy **MapReduce** logic.
- Standardized **CI/CD** for data pipelines with **AWS CodePipeline**, CodeBuild, **Jenkins**, **Git**, and **Terraform**-managed infrastructure across dev, test, and production.
- Owned migration programs end to end, from source profiling and gap analysis through parallel-run validation, cutover planning, and post-migration reconciliation.
- Enforced **AWS IAM** role-based access, Secrets Manager credential rotation, and encryption controls on protected financial, government, and healthcare datasets.
- Lowered cloud spend through S3 storage tiering, EMR transient and spot capacity, **Redshift** concurrency scaling, and **Snowflake** warehouse auto-suspend.
- Automated secure file handling, job scheduling, archival, operational checks, and restart processing with UNIX, Shell, and **PowerShell** scripting.
- Owned disaster recovery and business continuity procedures including scheduled rollback testing, and supported hybrid connectivity between on-premises systems and cloud platforms.
- Anchored production support and on-call rotation with root-cause analysis and permanent fixes, and mentored mid-level and associate engineers on **Spark** tuning and pipeline design.
- Governed data lineage, cataloguing, and master data management across regulated domains, aligning pipeline design with HIPAA, PII, and public-sector retention requirements.
- Translated business requirements into source-to-target mapping documents, technical design specifications, and testable acceptance criteria for downstream QA and audit teams.
- Partnered with architects, QA, DevOps, and business stakeholders across the full SDLC in Agile and Scrum delivery, including sprint and PI planning.

CORE COMPETENCIES

Cloud Data Architecture • ETL/ELT Pipeline Development • Data Migration & Modernization • Big Data Processing • Data Warehousing • Dimensional Modeling • Data Lake Implementation • Batch & Real-Time Streaming • Data Integration • Data Quality & Reconciliation • Performance Tuning & Optimization • Workflow Orchestration • CI/CD Automation • Data Governance • Master Data Management • Source-to-Target Mapping • LLM & Agent Workflow Engineering • Production Support • Agile/Scrum Delivery

TECHNICAL SKILLS

Core Stack	Python, PySpark, Spark SQL, SQL, AWS Glue, Amazon S3, Amazon Redshift, Amazon EMR, AWS Lambda, AWS Step Functions, Snowflake, Apache Airflow, Amazon MWAA
AWS	AWS Glue, Amazon S3 Data Lake, Amazon Redshift, Amazon RDS, Amazon Athena, Amazon DynamoDB, Amazon EMR, AWS Lambda, AWS Step Functions, Amazon Kinesis, Amazon CloudWatch, AWS IAM, AWS Secrets Manager, Amazon MWAA, AWS Lake Formation, Amazon SQS, Amazon SNS, Amazon Aurora, Amazon API Gateway, Amazon EventBridge, AWS CodePipeline, AWS CodeBuild, Amazon Bedrock
Azure	Azure Data Factory, Azure Databricks, Delta Live Tables, Unity Catalog, Azure Synapse Analytics, Azure Data Lake Storage Gen2, Azure Blob Storage, Azure SQL Database, Azure Functions, Azure Event Hubs, Azure Stream Analytics, Azure Key Vault, Azure Monitor, Azure DevOps, Azure Entra ID, Cosmos DB, Azure Purview
Google Cloud	BigQuery, Cloud Dataproc, Cloud Dataflow, Cloud Composer, Cloud Storage, Cloud Pub/Sub, Cloud SQL, Cloud Functions, Cloud Build, Cloud IAM, Secret Manager, Dataplex, Looker
Microsoft Fabric	OneLake, Lakehouse, Fabric Data Factory, Dataflows Gen2, Synapse Data Engineering, Synapse Data Warehouse, Fabric Notebooks, Real-Time Intelligence, Eventhouse / KQL, Power BI Direct Lake, Shortcuts, Data Activator, Fabric Capacity (F-SKU), Microsoft Purview
AI / LLM	LangGraph, LangChain, Anthropic & OpenAI APIs, Prompt Engineering, Structured JSON Outputs, Agent Tool Calling, Retry & Feedback Loops, Output Guardrails, Schema Validation, Confidence Thresholds, Human-in-the-Loop Review, Vector Search / RAG
Languages	Python, PySpark, Scala, SQL, PL/SQL, Java, Shell Script, PowerShell, Perl Script
Databases	Oracle, MySQL, SQL Server, PostgreSQL, MongoDB, DynamoDB, Cosmos DB, Cassandra, Snowflake, Teradata, Teradata BTEQ, Amazon Aurora
Big Data	Spark, Spark SQL, Spark Streaming, Databricks, Delta Lake, Delta Live Tables (DLT), DLT Expectations, Unity Catalog, Apache Iceberg, Hive, Kafka, Sqoop, Oozie, Zookeeper, Flume, Impala, Storm, Solr, Pig, Druid, Hadoop, HDFS, MapReduce, YARN, Cloudera, Hortonworks, MapR, Databricks
ETL & Warehousing	Informatica PowerCenter, SSIS, SSAS, SSRS, dbt, ETL/ELT Development, Data Modeling, Star Schema, Snowflake Schema, Slowly Changing Dimensions, OLTP, OLAP, Data Marts, ODS, Erwin, Medallion Architecture
DevOps & CI/CD	Apache Airflow, Amazon MWAA, Terraform, Docker, Kubernetes, Jenkins, Oozie, UC4 / Automic, Grafana, AWS CodePipeline, AWS CodeBuild, Azure DevOps, Git, GitHub, SVN
BI & Practices	Amazon QuickSight, Tableau, Power BI, Looker, Agile, Scrum, SDLC, Data Governance, Data Lineage, Data Quality Assurance, Master Data Management, Jira, Confluence, ServiceNow, Windows, UNIX, Linux

CERTIFICATIONS

- Microsoft Certified: Fabric Analytics Engineer Associate (DP-600)
- Microsoft Certified: Power BI Data Analyst Associate (PL-300)
- Microsoft Certified: Azure Data Engineer Associate (DP-203)
- Databricks Certified Data Engineer Professional
- Snowflake SnowPro Certification
- AWS Certified AI Practitioner

PROFESSIONAL EXPERIENCE

HighRadius Corporation Houston, TX	Nov 2024 – Present
Senior Data Engineer	
- Led a team of 5 engineers across the receivables data platform, owning the delivery roadmap, code review standards, and production readiness for 60+ pipelines. - Cut nightly receivables pipeline runtime from 4.5 hours to 38 minutes by tuning Spark partitioning, broadcast joins, and executor memory allocation across 60+ AWS Glue jobs. - Reduced cash-application exceptions by 34% by automating data-quality checks on invoice totals, payment matching, and source-to-target record counts across 45 daily finance feeds. - Consolidated 18 ERP, CRM, bank, and API source feeds into a single curated Amazon Redshift and Snowflake layer, shortening month-end close reporting from 6 days to 2.	

- Raised on-time pipeline delivery against SLA from **91%** to **99.2%** by orchestrating **40+** multi-stage finance workflows in **Amazon MWAA** with dependency, retry, and alerting logic.
- Handled **12M** payment and remittance events per day through **Amazon Kinesis**, **Kafka**, and **Spark Streaming**, reducing downstream data latency from **4 hours** to under **5 minutes**.
- Trimmed manual analyst review on deduction coding and remittance classification from **12 minutes** to **90 seconds** per item by building **LangGraph** agent workflows with tool-calling over internal finance APIs, covering **25K** items per month.
- Lifted field-extraction accuracy on unstructured remittance documents from **78%** to **96%** by engineering structured prompts with few-shot examples and enforced JSON schema outputs against LLM APIs.
- Held hallucinated field values under **0.5%** across **25K** monthly inferences by implementing guardrails for schema validation, confidence thresholds, bounded retry loops, and human-in-the-loop escalation.
- Administered the **Microsoft Fabric** tenant and workspace architecture, establishing dev, test, and production separation with deployment pipelines and infrastructure-as-code.
- Managed Fabric capacity across **F-SKU** tiers, right-sizing compute against workload profiles and reducing capacity spend by **22%** without breaching refresh SLAs.
- Shortened finance reporting refresh from **45 minutes** to near-real-time for **200+** business users by building a **Microsoft Fabric** lakehouse on **OneLake** with **Dataflows Gen2** and **Power BI Direct Lake** semantic models.
- Layered Bronze, Silver, and Gold medallion zones across 9 data domains with schema validation, cleansing, normalization, and reconciliation testing at each boundary.
- Owned disaster recovery and business continuity procedures for the data platform, running scheduled rollback tests and restoring service within a 4-hour RTO.
- Shortened release cycle time from **3 days** to **4 hours** by establishing **CI/CD** in **AWS CodePipeline**, **CodeBuild**, and **Jenkins** with automated testing across 15 pipeline repositories.
- Completed requirements analysis, application design, development, testing, and production operations for financial SaaS data solutions serving invoicing, payments, remittances, deductions, and collections.
- Developed DDL, DML, views, stored procedures, functions, triggers, and packages supporting financial staging, transformation, and reporting layers.
- Designed custom **AWS Lambda**, **Amazon EMR**, and **Databricks** components for cleansing, standardization, validation, and exception processing of high-volume transaction data.
- Created **Airflow** operators and parameterized workflows supporting daily, intraday, month-end, and backfill processing cycles.
- Implemented **Spark SQL** and **PySpark** transformations on **Amazon EMR** and **Databricks** for receivables analytics, cash application, and financial reporting.
- Monitored scheduled workflows through **AWS Step Functions**, **Amazon CloudWatch**, and **Grafana** dashboards, investigating failures and restoring processing within delivery SLAs.
- Decoupled finance event processing with Amazon **SQS**, **SNS**, and **EventBridge**, fanning **12M** daily payment events out to downstream consumers with dead-letter queues and replay for failed deliveries.
- Exposed curated finance datasets through Amazon **API Gateway** with Lambda authorizers, throttling, and usage plans for **20+** downstream applications.
- Containerized **PySpark** and **Python** transformation jobs with **Docker** and deployed them on **Kubernetes**, standardizing runtime dependencies across dev, test, and production.
- Shipped **AWS Lambda** functions ingesting finance data from Amazon **Aurora**, **Amazon RDS**, REST APIs, secure file transfers, and event-driven sources.
- Built **PL/SQL** utilities automating reconciliation, cleansing, transformation, and validation of financial transactions at scale.
- Applied **Informatica** Data Validation Option to compare source and target results and confirm financial-data accuracy across migrated platforms.
- Architected Shell and UNIX scripts for secure file handling, job automation, archival, operational checks, and restart processing.
- Managed epics, stories, defects, and QA validation in **Jira**, and participated in Agile ceremonies, release planning, retrospectives, and PI planning across distributed teams.

Environment: AWS Glue, AWS Lambda, AWS Step Functions, Amazon MWAA, Apache Airflow, Amazon EMR, Databricks, Amazon Redshift, Snowflake, Amazon Kinesis, Kafka, Spark Streaming, PySpark, Spark SQL, Python, PL/SQL, MySQL, Cassandra, LangGraph, LangChain, Anthropic & OpenAI APIs, Amazon Bedrock, Microsoft Fabric, OneLake, Lakehouse, Dataflows Gen2, Fabric Deployment Pipelines, Fabric Capacity (F-SKU), Power BI Direct Lake, Medallion Architecture, Amazon SQS, Amazon SNS, Amazon EventBridge, Amazon API Gateway, Amazon Aurora, Docker, Kubernetes, Grafana, AWS CodePipeline, CodeBuild, Jenkins, Git, Jira, Amazon QuickSight.

State of Illinois | Springfield, IL

Jan 2023 – Oct 2024

GCP Data Engineer

- Owned the cloud migration roadmap for 3 agency programs, presenting architecture and cutover plans to the data governance board for approval.
- Migrated 45 on-premises **Oracle** ETL workloads to **Cloud Dataproc**, **BigQuery**, and **Cloud Storage**, reducing annual infrastructure spend by **\$310K** while preserving source-to-target reconciliation controls.
- Decreased agency reporting runtime from **3.5 hours** to **22 minutes** by redesigning **BigQuery** partitioning, clustering, and join strategy across **12TB** of program data.

- Ingested **18TB** of structured and semi-structured agency data into **Cloud Storage** and **BigQuery**, enabling 9 public-service reporting products used by **400+** agency staff.
- Secured 25 protected government datasets by implementing **Cloud IAM** role-based access and **Secret Manager** credential rotation, closing 14 audit findings with zero repeat exceptions.
- Eliminated **10 hours** of manual deployment effort per week by automating build and release for **Python**, **SQL**, and **Spark** components through **Cloud Build**, Cloud Source Repositories, and **Jenkins**.
- Rolled out 22 **Hive** UDFs applying policy-driven business rules consistently across programs and reporting periods, cutting per-cycle rework by **15 hours**.
- Deployed public-sector data solutions across **Hadoop**, **Cloudera**, **Spark**, and **Google Cloud** services supporting agency operational and reporting workloads.
- Produced ingestion and orchestration workflows with **Cloud Dataflow**, **Cloud Dataproc**, **Databricks**, and **Cloud Composer** for scheduled agency-data processing.
- Modernized legacy **SQL** databases and file-based feeds into **Cloud Storage**, **Cloud SQL**, and **BigQuery** warehouse layers.
- Used bulk-load and warehouse-native transfer patterns to move data efficiently between **Cloud Storage**, **BigQuery**, and **Cloud Dataproc**.
- Authored batch and streaming solutions with **Spark Streaming**, **Cloud Pub/Sub**, **Kafka**, and schema-aware processing frameworks.
- Processed records across multiple data formats using **Scala**, **PySpark**, **Spark SQL**, and reusable transformation libraries.
- Employed Apache **Sqoop** to migrate recurring **MySQL** and relational extracts into **HDFS** for legacy downstream processing.
- Adopted **PySpark** RDDs, DataFrames, and **Spark SQL** for data profiling, cleansing, standardization, and aggregation across agency datasets.
- Established **Spark Streaming** micro-batches for time-sensitive feeds feeding downstream batch processing.
- Optimized **PySpark** and **Spark SQL** jobs through partitioning, caching, join strategy, and configuration tuning.
- Administered **Git** repositories, pull-request workflows, branching standards, and versioned deployment artifacts.

Environment: Google Cloud Platform, BigQuery, Cloud Dataproc, Cloud Dataflow, Cloud Composer, Cloud Storage, Cloud Pub/Sub, Cloud SQL, Cloud IAM, Secret Manager, Cloud Build, Cloud Functions, Looker, Databricks, Hadoop, Cloudera, Hortonworks, Hive, Sqoop, Spark Streaming, Kafka, Scala, PySpark, Python, PowerShell, Oracle, MySQL, Jenkins, Git, Jira.

InnovaCare Health | Athens, GA

Oct 2020 – Dec 2022

Data Engineer | AWS & Azure

- Served as technical lead for the hybrid AWS and **Azure** healthcare data platform, setting design standards and reviewing all pipeline changes before release.
- Built **AWS Glue** and **Azure Data Factory** ingestion for member, provider, eligibility, claims, and clinical-quality data across 14 source systems, processing **8M** claims records per month.
- Compressed the quality-measure reporting cycle from **3 weeks** to **4 days** by delivering **PySpark** transformations on **Amazon EMR** and **Azure Databricks**.
- Built **Delta Live Tables** pipelines with **DLT** expectations (expect_or_drop and expect_or_fail) to enforce claims and eligibility data quality at ingestion, routing violations to quarantine rather than failing full runs.
- Governed access through **Unity Catalog** with catalog, schema, and table-level grants, column masking on PHI fields, and lineage tracking across the healthcare lakehouse.
- Curbed claims data defects by **41%** by adding validation, reconciliation, and exception-handling controls to 30 production pipelines.
- Lowered cluster compute cost by **28%** by right-sizing **Amazon EMR** and **Azure Databricks** capacity and shifting 35 batch jobs to transient clusters.
- Documented source-to-target mappings and hybrid AWS and **Azure** architecture for 20 healthcare data flows, cutting new-engineer onboarding from **6 weeks** to 2.
- Orchestrated upstream-to-downstream healthcare workflows across **AWS Glue** and **Azure Data Factory** pipelines using dependency controls, retries, and operational alerting.
- Automated scheduled, event-driven, and recurring jobs for daily claims, membership, provider, and quality-measure processing.
- Harnessed **Amazon DynamoDB** and **Azure Cosmos DB** for reference, catalog, and event data supporting member and provider processing workflows.
- Assembled high-level and application design documents covering ingestion, transformation, security, lineage, recovery, and downstream interfaces.
- Prepared mapping specifications and data-flow rules for claims, eligibility, provider, member, and clinical-quality datasets.
- Handled claims and event data in micro-batches with **Spark Streaming** and **Azure Event Hubs**, applying RDD and DataFrame transformations for timely analytics.
- Provisioned **Amazon EMR** and **Azure Databricks** clusters for batch and continuous processing with **AWS IAM**, Secrets Manager, and **Azure Key Vault** controls on protected health data.
- Governed the healthcare data lake with **AWS Lake Formation**, applying row- and column-level permissions and centralized catalog access across PHI datasets.
- Stood up **AWS Glue** pipelines for copy, filter, conditional, transformation, and **Spark**-based healthcare processing patterns.
- Operationalized build and release definitions with **AWS CodePipeline**, CodeBuild, **Azure DevOps**, and **Jenkins** for CI, testing, and controlled deployment.

- Landed curated healthcare datasets into **Amazon S3** and **Azure Data Lake Storage Gen2**, serving downstream analytics through **Amazon Redshift** and **Azure Synapse Analytics**.
- Specified **SQL**, stored procedures, triggers, and packages for healthcare staging, audit, transformation, and reporting databases.

Environment: AWS Glue, Amazon EMR, AWS Lambda, Amazon DynamoDB, Amazon RDS, Amazon Redshift, Amazon S3, AWS Lake Formation, AWS IAM, AWS Secrets Manager, Azure Data Factory, Azure Databricks, Delta Live Tables, Unity Catalog, Azure Synapse Analytics, Azure Data Lake Storage Gen2, Azure Event Hubs, Azure Key Vault, Azure DevOps, Cosmos DB, PySpark, Spark Streaming, Spark SQL, Python, Scala, Hive, Oozie, Oracle 11g, MySQL, Snowflake, Maven, Linux, AWS CodePipeline, CodeBuild, Jenkins.

Giant Food | McLean, VA

Apr 2018 – Sep 2020

Big Data Developer

- Built **Spark** applications in **Scala** and **PySpark** processing **9M** retail transactions daily, replacing 45 legacy **MapReduce** jobs and cutting the nightly batch window by **3 hours**.
- Operated **Hadoop** workloads across **MapR** and **Cloudera** distributions, tuning cluster configuration and resource allocation for high-volume retail transaction processing.
- Compressed merchandising query response from **6 minutes** to **25 seconds** by migrating 30 recurring analyses from **Hive** to **Impala**.
- Improved daily POS and inventory feed throughput by **38%** by optimizing **MapReduce** jobs and compression strategy across 160 store locations.
- Removed **12 hours** of weekly manual scheduling by automating dependent **Hive** and Pig jobs through **Oozie** workflows.
- Coded UDFs applying retail business rules for product hierarchy, store, promotion, and transaction processing.
- Exported curated retail datasets to relational databases supporting BI dashboards, merchandising analysis, and operational reporting.
- Wrote Java **MapReduce** programs to extract, transform, and aggregate XML, JSON, CSV, and compressed retail files.

Environment: Hadoop, MapReduce, HDFS, Hive, Pig, Spark, Impala, HBase, Java, SQL, Cloudera Manager, MapR, Scala, Sqoop, Storm, Solr, Mahout, Flume, Oozie, Eclipse.

Morningstar, Inc. | Chicago, IL

Jun 2016 – Mar 2018

ETL Developer

- Met market-data delivery windows **99.4%** of the time across 55 investment feeds by applying **Informatica** pushdown optimization and parallel workflow execution, reducing run time by **40%**.
- Supplied SCD Type 1 and Type 2 star schemas supporting 25 investment research products.
- Scaled back UAT defects by **32%** through systematic source-to-target reconciliation testing across security, fund, pricing, and portfolio datasets.
- Analyzed investment-data requirements from data architects and translated them into source-to-target ETL specifications.
- Mapped out **Informatica** mappings, sessions, workflows, UNIX scripts, and parameter files to process investment and market data.
- Scheduled and monitored production ETL workflows through **UC4/Automic**, coordinating dependencies across 55 investment data feeds and time-critical delivery windows.
- Provided 24x7 on-call support for production ETL workflows and critical market-data deliveries.

Environment: Informatica PowerCenter, UC4 / Automic, Oracle, UNIX Shell Script, SQL, Star Schema, Slowly Changing Dimensions, Windows.

Transamerica | Baltimore, MD

Jan 2014 – May 2016

ETL Developer

- Loaded **22M** policy, premium, and claims records into the Operational Data Store via **Informatica PowerCenter** mappings at **99.8%** reconciliation accuracy.
- Handed over 15 **Erwin** logical and physical data models supporting relational and dimensional insurance reporting databases.
- Drove down recurring reporting effort by **60 hours** per month by building **SSIS** packages and **SSRS** reports for policy, customer, and financial data.
- Instituted enterprise data-warehouse solutions for insurance, retirement, policy, premium, claims, customer, and financial reporting data.
- Drafted mappings with joiners, lookups, expressions, filters, routers, aggregators, sorters, stored procedures, and update strategies.
- Wrote Teradata **BTEQ** scripts for bulk extract, load, and validation of policy and premium data across warehouse environments.
- Composed **SSIS** packages loading **SQL Server** and **SSAS** structures for insurance and retirement reporting.
- Delivered sequential and concurrent **Informatica** sessions for reliable execution of dependent mappings.

Environment: Informatica PowerCenter 8.6.1, SQL Server 2005, SSIS, SSRS, SSAS, Oracle, Teradata, BTEQ, Erwin, VBA, RDBMS, FTP, SFTP.

EDUCATION

Master's in Information Technology | Catholic University of America | Washington, DC

Aug 2012 – Dec 2013

Bachelor of Computer Science and Technology | Lovely Professional University | India

Aug 2008 – Apr 2012