

Service Name

Amazon Data Firehose

Service Overview

Amazon Data Firehose is a fully managed service for delivering real-time streaming data to destinations such as Amazon S3, Amazon Redshift, Amazon OpenSearch Service, Splunk, Apache Iceberg Tables, and custom HTTP endpoints. The service automatically buffers, transforms, compresses, and delivers streaming data without requiring customers to write applications or manage resources. It supports data ingestion from multiple sources including Amazon Kinesis Data Streams, Amazon MSK, Direct PUT API, and AWS services like CloudWatch Logs. Data Firehose automatically scales to match the data throughput and provides near real-time data delivery with configurable buffer sizes and intervals.

Key Use Cases

- Streaming data into data lakes and warehouses: Loading real-time data from applications into Amazon S3 for analytics and long-term storage
- Security monitoring and threat detection: Streaming security logs and events to security analytics platforms like Splunk for real-time monitoring
- Machine learning data pipelines: Feeding streaming data into ML platforms and analytics services for real-time inference and model training
- Business intelligence and analytics: Delivering streaming business metrics to data warehouses like Amazon Redshift for real-time dashboards
- Log aggregation and analysis: Centralizing application and infrastructure logs for monitoring, troubleshooting, and compliance

Features

- **Fully Managed Service:** No infrastructure management required, automatic scaling and provisioning
- **Multiple Data Sources:** Supports Kinesis Data Streams, Amazon MSK, Direct PUT, and AWS service logs
- **Data Transformation:** Real-time data transformation using AWS Lambda functions before delivery
- **Format Conversion:** Convert data to columnar formats like Apache Parquet and ORC for optimized analytics
- **Dynamic Partitioning:** Automatically partition data in S3 based on content for improved query performance
- **Multiple Destinations:** Deliver to Amazon S3, Redshift, OpenSearch, Splunk, Snowflake, and HTTP endpoints

- **Buffer Management:** Configurable buffer sizes and intervals for optimized delivery batches
- **Encryption and Security:** Server-side encryption with AWS KMS and IAM-based access control
- **Error Handling:** Automatic retry logic and error logging with failed record backup
- **Monitoring Integration:** CloudWatch metrics and logging for stream monitoring and troubleshooting

Value Proposition

Amazon Data Firehose eliminates the complexity of building and managing streaming data pipelines by providing a fully managed, serverless solution. Customers choose Firehose for its zero infrastructure management, automatic scaling, and pay-as-you-go pricing model with no upfront commitments:

- The service reduces time-to-value by enabling immediate data streaming without writing custom applications or provisioning resources
- It offers built-in reliability features including automatic retries, error handling, and data backup capabilities
- Integration with over 30 AWS services and third-party destinations provides flexibility for diverse analytics architectures
- The service supports real-time data transformations and format conversions, optimizing data for downstream analytics while maintaining high availability and durability

Service Integration

Amazon Data Firehose integrates deeply with core AWS services across the data analytics ecosystem. It connects with Amazon Kinesis Data Streams and Amazon MSK as streaming data sources, while supporting direct ingestion from AWS services like CloudWatch Logs, VPC Flow Logs, and AWS WAF:

- For data processing, it integrates with AWS Lambda for real-time transformations and AWS Glue for metadata management and schema discovery
- Destination integrations include Amazon S3 for data lakes, Amazon Redshift for data warehousing, and Amazon OpenSearch Service for search analytics
- Security integrations include AWS KMS for encryption and AWS IAM for access control
- Monitoring is provided through Amazon CloudWatch for metrics and logging, while AWS Secrets Manager handles credential management for external destinations

Onboarding and Offboarding

Customers get started with Amazon Data Firehose by creating a delivery stream through the AWS Management Console, CLI, or SDKs. The process involves selecting a data source (Direct PUT, Kinesis Data Streams, or Amazon MSK), configuring optional data transformations, choosing a destination, and setting buffer configurations:

- Prerequisites include signing up for an AWS account and creating appropriate IAM roles for service permissions
- The service provides tutorials and sample code for common integration patterns
- Getting started is streamlined through the console wizard that guides users through stream creation, transformation setup, and destination configuration
- No infrastructure provisioning or application development is required, allowing customers to begin streaming data within minutes of setup

Getting Started Guide: <https://docs.aws.amazon.com/firehose/latest/dev/what-is-this-service.html>

Data Removal

Offboarding from Amazon Data Firehose involves stopping data ingestion, allowing existing buffered data to be delivered, and then deleting the delivery stream. Customers can pause data delivery temporarily or permanently delete streams through the console or APIs. Data removal depends on destination configurations - for S3 destinations, customers must manually delete objects from buckets; for other destinations like Redshift or OpenSearch, data persists according to those services' retention policies. Source data backup in S3 requires separate cleanup. Failed records stored in error buckets also need manual deletion during offboarding.

Backup/Restore and Disaster Recovery

Amazon Data Firehose provides built-in backup capabilities by optionally storing source data in Amazon S3 concurrently with destination delivery. This includes failed record backup for troubleshooting and source record backup for data transformation scenarios. The service offers at-least-once delivery semantics with automatic retry logic for up to 24 hours for DirectPut sources. For disaster recovery, Firehose relies on underlying AWS service durability and cross-region replication capabilities of destinations like S3. The service itself is highly available within regions but does not provide native cross-region replication.

Backup Capabilities

Amazon Data Firehose includes automatic backup features for data reliability and troubleshooting. Source data backup is available when using data transformations, storing original records in a separate S3 bucket:

- Failed record backup automatically saves records that couldn't be delivered to destinations in designated S3 error folders
- The service buffers data for up to 24 hours during delivery failures for DirectPut sources
- However, Firehose does not integrate directly with AWS Backup service as it's primarily a data streaming rather than storage service

Disaster Recovery

Amazon Data Firehose does not directly integrate with AWS Elastic Disaster Recovery as it's a data streaming service rather than compute infrastructure. The service provides high availability within AWS regions through managed infrastructure and automatic failover capabilities:

- Disaster recovery strategies typically involve replicating destination data stores (like S3 buckets or Redshift clusters) to other regions rather than replicating the Firehose streams themselves
- Cross-region disaster recovery is achieved through destination-level replication and recreating Firehose streams in alternate regions when needed

Recovery Objectives

Amazon Data Firehose helps customers meet RPO objectives through its reliable data delivery with at-least-once semantics and automatic retry mechanisms. RPO can be as low as seconds due to real-time streaming capabilities and configurable buffer intervals. RTO objectives are supported through automatic scaling, managed infrastructure, and quick stream recreation capabilities. The service's 99.9% SLA uptime guarantee and built-in error handling help maintain consistent data flow. Buffer management allows customers to balance between data freshness and delivery efficiency to meet specific RPO requirements.

Service Constraints

Service Quotas: <https://docs.aws.amazon.com/firehose/latest/dev/limits.html>

FAQ: <https://aws.amazon.com/firehose/faqs/>

Technical Requirements

Service Documentation: <https://docs.aws.amazon.com/firehose/latest/dev/index.html>

User Guide: <https://docs.aws.amazon.com/firehose/latest/dev/what-is-this-service.html>

API Reference:

<https://docs.aws.amazon.com/firehose/latest/APIReference/Welcome.html>

Pricing Overview

Please see the AWS UK G Cloud 15 Pricing Document accompanying this service in the Digital Marketplace.

Public Pricing: <https://aws.amazon.com/firehose/pricing/>

Cost Optimization

Amazon Data Firehose offers several unique cost optimization opportunities through its tiered pricing model and pay-as-you-go structure. Optimizing record sizes to avoid the 5KB minimum billing increment for Direct PUT sources can significantly reduce costs:

- Batch processing using PutRecordBatch API is more cost-effective than individual PutRecord calls
- Buffer size and interval tuning helps optimize delivery batches, reducing per-request costs at destinations
- Format conversion to columnar formats like Parquet reduces storage costs and improves query performance
- Dynamic partitioning, while adding processing costs, can reduce downstream query costs significantly
- Compression algorithms reduce data transfer and storage costs
- Choosing appropriate data sources (Direct PUT vs
- MSK pricing tiers) based on volume can optimize ingestion costs

Savings Considerations

Primary cost savings come from eliminating infrastructure management overhead and reducing operational complexity compared to self-managed streaming solutions. No upfront costs or minimum commitments provide flexibility for variable workloads:

- Tiered pricing rewards high-volume users with lower per-GB rates as usage scales
- Automatic scaling eliminates over-provisioning costs common with traditional infrastructure
- Built-in compression and format conversion reduce downstream storage and compute costs
- Error handling and retry mechanisms reduce data loss costs and operational overhead
- Integration with managed AWS services eliminates need for custom integration development
- The serverless model means customers only pay for actual data processing rather than idle infrastructure capacity, providing significant savings for variable or unpredictable workloads