



Self-Optimizing Cloud Data Pipelines via Machine Learning

Sophie Dubois
Asst Professor

Abstract

Cloud Data Pipelines enable enterprises to readily ingest, process, clean, and store large amounts of structured and unstructured data in cloud environments to drive analytics, business intelligence, and data-science workloads. However, designing and implementing such pipelines is non-trivial and challenging. Pipelines should be optimized for cost, speed, or any combination of the two but these objectives are at odds with each other. A data pipeline architecture that enables easy prototyping of data ingestion and transformation processes within any cloud platform is presented. Machine Learning (ML) is employed to inform scheduling and resource allocation decisions in order to reduce operational cost while ensuring acceptable latencies.

The objectives of optimizing Ingest, Transformation, and Enhanced ETL cloud data pipelines in real-time for cost and latency are accomplished. The four cloud providers—Google, Amazon, Microsoft, and IBM—are supported, with data volumes ranging from a few megabytes to generating several gigabytes. Latency from minutes to hours can be supported without breaking the bank. ML models inform autoscaling groups, transformation resources, and scheduling. Cross-cloud portability through modular code-based connection-management further optimizes the development phases while improving code quality.

Keywords: Cloud Data Pipelines, Real-Time ETL Optimization, ML-Driven Resource Scheduling, Cost–Latency Trade-off, Cross-Cloud Data Architecture, Intelligent Autoscaling, Data Ingestion Frameworks, Transformation Pipelines, Enhanced ETL, Cloud Portability, Modular Connection Management, Multi-Cloud Support (AWS, GCP, Azure, IBM), ML-Based Workload Scheduling, Resource Allocation Optimization, Data Engineering Automation, Operational Cost Reduction, Low-Latency Analytics Pipelines, Structured and Unstructured Data Processing, Cloud-Native Data Platforms, Production-Grade Data Pipelines.

1. Introduction

The costs and complexities of executing cloud workload make optimizing resource management a priority. Classical approaches utilize careful capacity provisioning through autoscaling, but controlling Quality of Service (QoS) in such systems often relies on best efforts.



QoS violations can introduce unexpected service downtimes, degrade user experiences, and decrease revenue. The portable, cloud-agnostic design of cloud data pipelines gives automated Machine Learning (ML) modellers a high-level view of the various influencing factors, offering data-driven solutions to resource allocation and scheduling problems. Tuning third-party services instead of application-managed services allows the use of less-knowledgeable yet easier-to-understand, can-be-executed-just-once ML models. Cloud data pipelines for batch workloads support scheduling decisions between service execution and asking an ML model to supply the execution decision. Public clouds provide data from many different users, which can be collected, labelled, and used for training, testing, or validating a wide range of ML models used to support many aspects of data pipeline operation.

The presented work studies self-optimizing cloud data pipelines using externalized ML models to support many aspects of QoS management. It defines the modeling and experimentation process, detailing the introduction of ML models for decision-making in resource allocation, scheduling, and scaling. The experiments involve data transfer or storage in multiple clouds, examples of data ingestion and data processing with CNNs or Web servers from a third-party service in multiple clouds. The cloud-in-cloud approach supports the Tests As A Service concept. The current scope focuses on ML models for scheduling and resource allocations, pursuing the creation and usage of a diverse set of models. Research work in the past few years and expected improvements by the emergence of new areas in the ML field drive model creation.

1.1. Overview of the Study's Objectives and Scope

In-cloud optimization of data pipelines for data integration and transformation has several objectives: minimizing cost without affecting latency; minimizing latency without degrading performance; and achieving high throughput at a minimal cost. These goals are pursued via machine learning models capable of predicting, for a given data workload, the best schedule that incurs the least cost for a required latency, or that achieves the best quality metric (cost or latency) for a given throughput demand. The research focuses on integration and transformation processes deployed on cloud data pipelines and hosted on the Google Cloud Platform. Other services (e.g. data storage, caching, and messaging) are assumed to be provided by the cloud

vendor management, and workloads are executed on large-scale cloud resources.

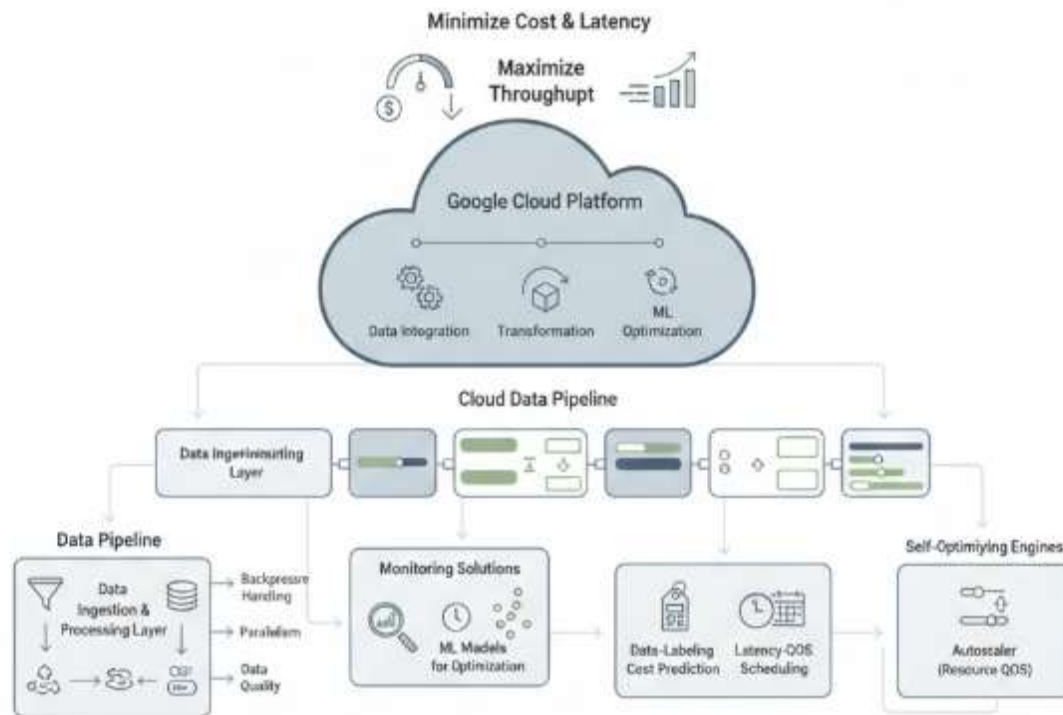


Fig 1: Autonomous Cloud Data Pipelines: A Machine Learning Framework for Multi-Objective Cost, Latency, and Throughput Optimization

Previous studies provided a modular architecture for data pipelines that run on any cloud platform. The architecture is implemented and the data ingestion and processing layer is described; it supports cloud-native streaming and batch services, handling backpressure, parallelism, and data quality. Monitoring solutions for the pipeline control plane collected data for training ML models that optimize data integration and transformation processes. Three kinds of ML-driven self-optimizing engines were designed and built. The first predicts how data-labeling requirements influence the overall pipeline cost; the second adjusts the scheduling of labels to minimize cost subject to fixed-latency QoS; and the third is a full-fledged autoscaler that controls pipeline resources based on user-defined QoS.

2. Background and Related Work



Data ingestion, integration, and transformation pipelines are a vital part of cloud workloads. They move data from multiple sources, such as web servers, application logs, third-party APIs, and monitoring systems, to storage solutions or data lakes designed for later analysis. Pipelines are often formed as a bouquet of cloud services offered by the main providers; these services cover data extraction from sources, transformation into the final format, and storage for further analysis. These flows consume resources and incur costs that can be tens to hundreds of dollars a day.

Thus, managing data flows should ideally be considered from a datacenter-wide perspective, in conjunction with other components such as batch jobs and machine learning training. However, these flows seldom get the attention they deserve. They are often coded as sets of ad-hoc scripts that get the job done for a specific need but without the guarantees of quality, maintainability, or reliability that one would expect from a production-grade system.

Existing systems take different approaches to optimizing data pipelines. Some of these systems leverage ML techniques to automatically make decisions, such as allocating resources or selecting the instance types of virtual machines that will run the data pipeline's components. Such optimizations may at times bring tangible gains, but they are not a panacea, since they focus on a single aspect of the flow and overlook others. A data pipeline can still benefit from automation and optimization even without the utilization of ML techniques—indeed, without ML techniques at all. The present work considers the entire data pipeline system, describing its architecture and implementation, and proposing key aspects of its operation that can be automated and optimized through ML techniques.

2.1. Related Research and Prior Art

Most existing data pipelines that consume cloud services focus the cost of data transfer and buffer storage while performing a simple best-fit scheduling of preprocessing tasks to minimize the use of cloud resources. A cost function is reduced by finding optimal transfer schedule to minimize the cost of data retrieval and data transfer. A smart queue service leverages the parallelism of cloud by redirecting the data stream from data producers to multiple data consumers based on the priority order and processing time. CloudCrane integrates multiple services of cloud infrastructure as a data segmentation framework which segments a data into several pieces and store into Cloud servers by retrieving the relevant credentials. Workload-aware scheduling considered multiple resource constraints such as time, budget, and data locality.



A framework called mCloud optimizes the job scheduling of multiple concurrent queries for data reservoir services provided by cloud service providers. With the considerations of monetary cost and Quality of Service (QoS) in job execution time, it assure that the fast response time with lower monetary cost. An end-to-end cloud-based framework is proposed to subsume the cost of data staging in pipeline pattern processing with reliable network. WhizCloud is a dynamic resource provisioning system for heterogeneous batch applications on clouds with dedicated resources. Instead of scheduling the resources for multiple users, it minimizes resource consumption for a single user who enforces batch jobs with very small temporal locality by utilizing the dedicated resources. Pipeline scheduling considered cost, delay and dependencies of cloud services by processing a specified request which includes known and unknown cloud services.

Equation 1) Pipeline performance model (latency, throughput, cost)

1.1 Notation (maps to the paper's concepts)

- Stages: $i \in \{1, \dots, N\}$
- Provisioned resources for stage i : r_i (e.g., number of workers/instances/slots)
- Arrival rate (workload intensity): λ (jobs/sec) — corresponds to workload pattern features
- Service rate per resource unit at stage i : μ_i (jobs/sec per unit)
- Total service capacity at stage i :

$$\mu_i^{\text{tot}}(r_i) = r_i \mu_i$$

1.2 Queueing-based latency per stage (step-by-step)

A common approximation for a saturated processing stage is an M/M/1 queue (or M/M/m; we keep it simple and interpretable).

Step 1 (stability condition):

For stage i to not explode in backlog,

$$\lambda < \mu_i^{\text{tot}}(r_i) = r_i \mu_i$$



Step 2 (utilization):

$$\rho_i = \frac{\lambda}{r_i \mu_i}$$

Step 3 (mean time in system for M/M/1):

$$W_i = \frac{1}{\mu_i^{\text{tot}}(r_i) - \lambda} = \frac{1}{r_i \mu_i - \lambda}$$

Step 4 (include fixed overheads):

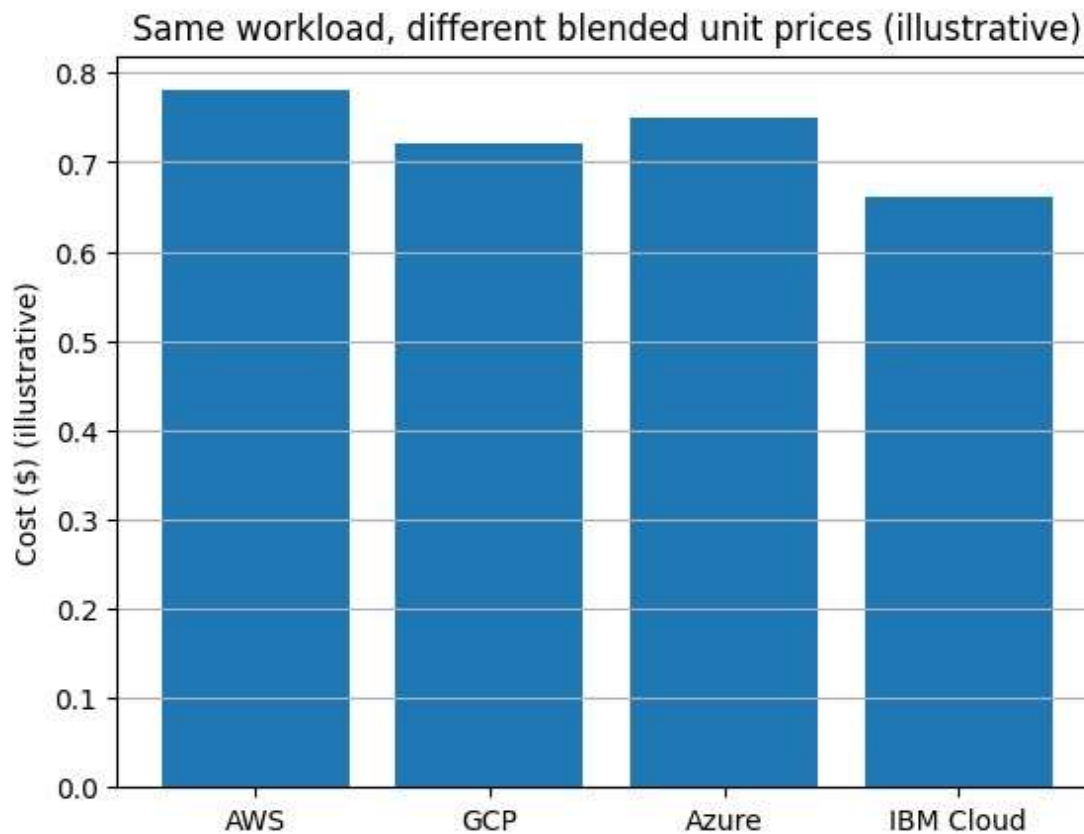
Real pipelines also have serialization, IO, orchestration, etc. Let that be d_i :

$$L_i(r_i) = W_i + d_i = \frac{1}{r_i \mu_i - \lambda} + d_i$$

Step 5 (end-to-end latency):

Summing stage latencies:

$$L_{\text{e2e}}(\mathbf{r}) = \sum_{i=1}^N L_i(r_i) = \sum_{i=1}^N \left(\frac{1}{r_i \mu_i - \lambda} + d_i \right)$$



3. Problem Formulation

An architect builds a house with wood. Later on, a carpenter fills the house with wood. This modularity allows specialization and portability to many regions with little effort. Data pipelines can be built upon similar modularity. Unit transactions must integrate, transform, store, and deliver data. The same pipeline is also usually used for a data warehouse load so scaling can be more efficient. Data always enters a data pipeline at the same point, but once it starts going through the house, it can take many paths. It may go through a streaming path or a batch path or some other process flow such as an ETL or ELT process for a data warehouse. During high-volume periods, backpressure can occur in the system, and the batch processing part of the system may have to run a few loads with a higher level of data latency to absorb the load. This part of the system will want to make sure it has available parallelism to assist with the necessary



data loads. Data pipelines rarely have any events triggered or controlled by the actual arrival of the data. A triggered event by the data arrival may push a follow-up alert or a warning into an event-driven system. A service mesh will be able to handle more fine-grained communication between all the microservice containers, but it has a heavier infrastructure load. Such setups may want to talk to the integration side during off-peak hours so the data handling can be done more seamlessly.

Cloud Engineering permits for the much-needed auto-scaling for microservice containers over the Ingress and Event Delivery Track. Reliability comes more from reliability riding over three edges or other topping patterns. But by becoming visible and computable, observability can also support emergent logic as the monitoring systems draw attention to latent patterns. Enhancing the ties among Schwartzian Analysis, SEO, SCHEMA, META-MODEL, and MASTER_MIND_TT is another future facet. All these other facets are generic add-ons to the architecture, allowing for different pipelining patterns while still allowing for the original situation where the domain-expert trained Data Quality Rules are the data pipeline flow controller.

Resource units (r)	Latency (sec) (illustrative)	Cost (\$) (illustrative)
1	128.0	0.12
2	68.0	0.24
3	48.0	0.36
4	38.0	0.48
5	32.0	0.6
6	28.0	0.72
7	25.14	0.84

3.1. Key Concepts in Cloud Data Pipeline Architecture

Data pipeline architecture in cloud computing environments consists of multiple integrated layers, each responsible for a specific functionality. The data ingress layer contains all the components and strategies needed to ingest any type of data from heterogeneous sources into the cloud provider. The data streaming layer enables the platform to react immediately to any event. The data storage layer offers cost-efficient storage solutions and services adapted to the data working set size. The data batch-processing layer executes the ELT or ETL process for



analytics workloads. Event-driven components allow processing on specific incoming data events without being triggered by user queries. Service meshes enforce and implement observability, reliability, and fault-tolerance primitives and rules for microservices deployed in the pipeline.

The modularity of the architecture design promotes isolation in the development cycle while supporting portability and compatibility with the services from different cloud providers. Each module can be implemented over the set of services offered by the cloud provider the organization has an agreement with. A GCP organization can use Pub/Sub and Dataflow for streaming data, whereas an AWS organization can use Kinesis and Lambda. For simpler data-pipeline use cases that do not need complex data processing, connector components to third-party providers, such as Fivetran and Striim, can be employed. Data-pipeline users are only required to satisfy the properties imposed by their choice of the modules.

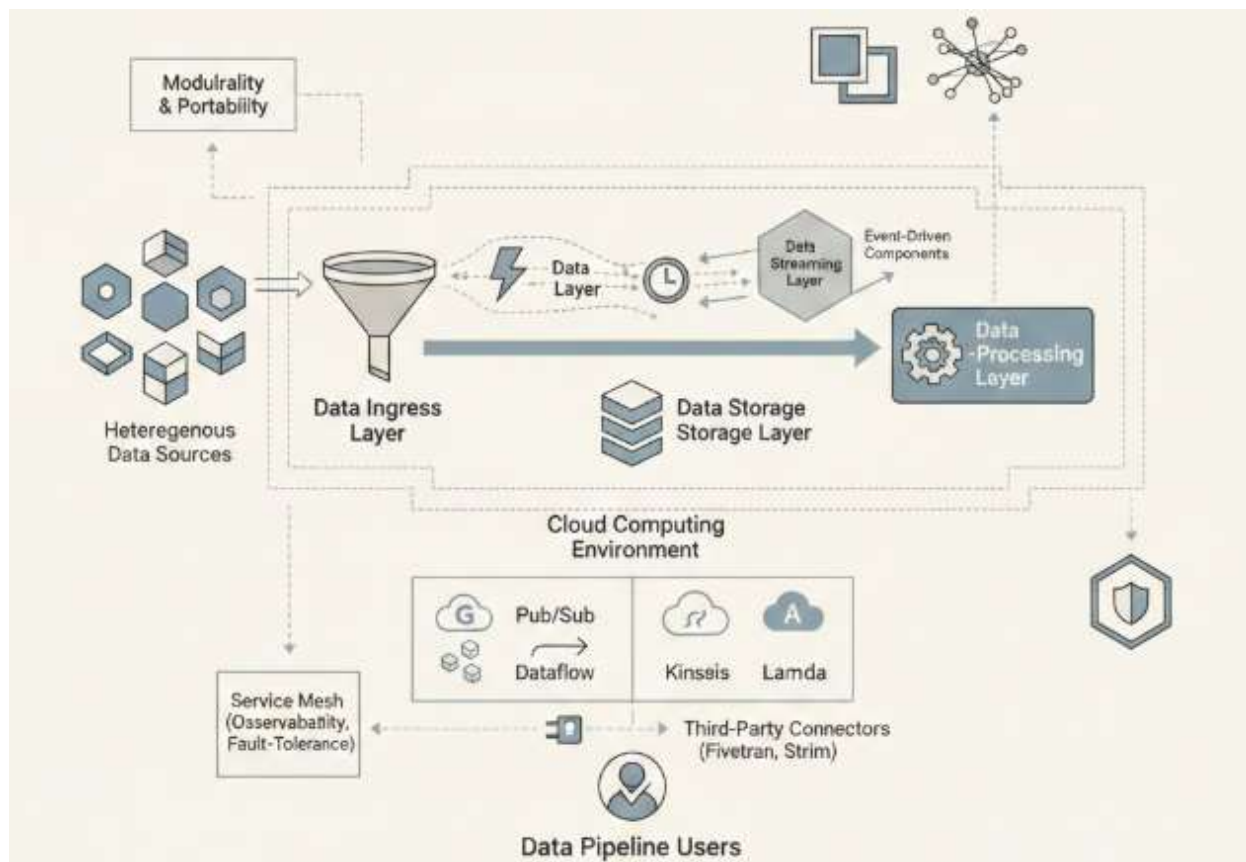


Fig 2: Poly-Cloud Architecture: A Modular Framework for Agnostic and Event-Driven Data Orchestration

4. Data Pipeline Architecture for Cloud Environments

The reference architecture of a data pipeline suitable for cloud deployment differentiates between four logical but closely interleaved layers: data ingress, processing, application, and control. Cloud systems are inherently multitenant, hence reliable and efficient separation of data is crucial. Each data-processing stage should flexibly span multiple cloud services, as the choice of service for a particular task may vary depending on the nature of the workload and the available cloud service at a given time. At the lowest layer, the control plane injects information,



policies, and control commands to the data-pipeline system, allowing it to optimally manage the services comprising the architecture and their interactions.

Such an architecture can serve as a skeleton for building cloud data pipelines: each logical layer can be instantiated using cloud services already in place, with the support of an orchestration system. It holds specific patterns for modularizing new data integration-and-transformation processes that can be made available as additional cloud services. Many additional capabilities are already managed natively by the cloud vendors but may require a particular design and implementation. Three of these capabilities are data quality, data security, and monitoring. A cloud data pipeline built according to this architecture can flexibly span multiple cloud-services providers whenever required and can be deployed and operated according to the principles of chaos engineering and data-driven SRE.

4.1. Data Integration and Transformation Processes

Data connectors enable efficient interaction with external data sources, while schema management systems support the auto-discovery of input and output data structures. Data quality rules detect anomalies in data generated by unreliable sources, whereas transformation pipelines enforce specific business rules on the data prior to utilization. Since data pipelines often rely on data from external partners, idempotency is an essential attribute for all pipeline components. By accumulating knowledge of external service responses, data pipelines can tolerate failures from those services without iteration.

Data pipelines sourced from external services should also exhibit fault tolerance. For example, change-data-capture components consuming Kafka streams from unreliable partners may skip events referencing unstable partitions or registers with inconsistent states. To divert initialization failures into an alternative control stream and attempt a new initialization attempt after N successful cycles in the main information topology, the originator must also keep track of the state from the last successful partition read. The usage of a service mesh transparently solves many of these issues and brings added observability and reliability primitives embedded into the architecture.



Equation 2) Cost model (what you optimize against latency)

2.1 Cost per stage

Let:

- p_i = price rate (\$ per resource-unit per second) for stage i
- T = time window (sec) you keep resources allocated

Then stage cost:

$$C_i(r_i) = p_i r_i T$$

Total cost:

$$C_{\text{tot}}(\mathbf{r}) = \sum_{i=1}^N p_i r_i T$$

4.2. Data Storage Solutions for Cloud Systems

Architectural variety offers a rich palette of data storage options, shaped by factors such as scale, dataset structure and access modes. The classic dichotomy of file systems (e.g., NTFS, HDFS) vs. databases (relational or NoSQL) becomes an oversimplification in the cloud. Object storage provides a S3-like interface to retrieve and store virtually unlimited volumes of unstructured data with low latency and reasonable economic costs; however, its lack of indexing and partitioning capabilities makes it unsuitable for large-scale ad-hoc analysis. Columnar storage (e.g., Parquet, ORC) overcomes these limitations and fuels parallel data warehouses, such as Redshift, BigQuery, and Synapse, in which many concurrent queries can each quickly scan a large dataset. Different consistency models fit various use cases and access patterns: read-after-write consistency for the latest data (e.g., social-media timelines); eventual consistency for highly-scalable read-heavy workloads (NoSQL systems); and configurable consistency between the two extremes (e.g., Amazon S3). Replication across datacenters supports disaster recovery and ultra-low-latency content delivery, while also exposing a cost-performance trade-off: latency and read throughput advantages for read-heavy workloads against increased expense for write-heavy loading patterns—read replicas can be stale, allowing lower storage costs per replica.



Layered storage provides a logical alternative: only the most recent copy is kept in hot S3 storage, while historical data with less frequent access is archived in warmer (e.g., S3-IA) and frozen (e.g., Glacier) storage tiers. Increasing the number of available layers enables a smoother cost-performance trade-off; however, a sparse dataset with only a few layers might not yield optimal cost savings with the same granularity.

Continued growth of fully-managed cloud data warehousing services like Amazon Redshift and Google BigQuery blurs the lines between online analytical processing (OLAP) systems and online transaction processing (OLTP) databases. Serverless architectures decouple resource provisioning from consumption, dynamically allocating resources to make ad-hoc queries responsive when using an analytic data warehouse. Reinforcement learning can be invoked to choose an appropriate cold storage layer throughout the pipeline, and the schedule can be combined with the number of layers to determine costs for data load-and-query use cases.

5. Machine Learning for Optimization

Machine learning (ML) techniques assist in determining configurations that result in suboptimal performance of data pipelines, such as tuning parameters, selecting an adequate pattern for scheduling component execution, and appropriate resources for data connector components. Training pipelines require collecting heterogeneous data from all layers of the architecture and modeling all processes controlled by ML algorithms or engines. Continuous monitoring and validation are essential to ensure model accuracy, and input data need to be periodically validated and relabeled before being retrained.

Features reflect the current workload pattern, the quality seals applied in upstream processes, resource consumption levels collected from monitoring solutions, data locality across cloud regions and availability zones, or historical information on the interaction of these characteristics with the QoS defined to execute specific data pipeline flows (for instance, latency, throughput, and adherence to service level agreement rules). Data collected from running the declared pipelines, such as resource utilization, progress of ongoing executions, and workload completion time, feed the subsystems in charge of making predictions or selecting the execution plans of the data flows. Based on these features, different ML models—such as reinforcement learning (RL), Markov decision processes, or supervised predictors—can be implemented to address the integration and processing of data pipelines, supporting predictive execution planning and responsive resource allocation. Model accuracy is determined by the modeling



purpose. Predictors are trained using historical information to classify the outcome argument, while RL and Markov decision processes are trained through interactions with the environment.

Cloud	Cost (\$) (illustrative)
AWS	0.78
GCP	0.72
Azure	0.75
IBM Cloud	0.66

5.1. Feature Engineering and Data Collection

Feature selection and customization are key elements of any supervised ML task. Although previous work on self-optimizing ML models has focused exclusively on using resource consumption and environment features (e.g., workload type and resource availability) to predict optimal values of parameters for client-driven services, a much broader set of features can be engineered to produce accurate predictions for pipeline scheduling and resource allocation. Such features can be grouped into the following four categories.

First, workload pattern features—such as incoming traffic rates, temporal traffic concentration, bursting patterns, and user/device locality—indicate how load will change over time and the balancing requirements across machines. The second group of features describes resource consumption or utilization for executors and data stores (long-term, short-term, or current values): CPU, memory, network, and disk throughput. The third group collects resource location information: the distance or transfer delay between clients and data stores or between two different data stores. Finally, quality-of-service (QoS) requirement features indicate the minimum acceptable latency/throughput values for a given job. Their unfulfillment can generate cost penalties or loss of service-level agreements.

Constructing a dataset for the model training, validation, and deployment phases requires using a tracer tool to record the model features, labeling, and resource usage when executing pipelines built as data flows, including data processing with Stream Processing Engines (SPEs) and Data Flow Engines (DFEs). The data from these pipelines can then be used to build, validate, and deploy ML models that predict the runtime cost (in monetary units or time) of pipeline components and resource schedules needed to satisfy QoS requirements. As an

additional option, the model can also be trained to predict the optimal values for parameters selected beforehand.

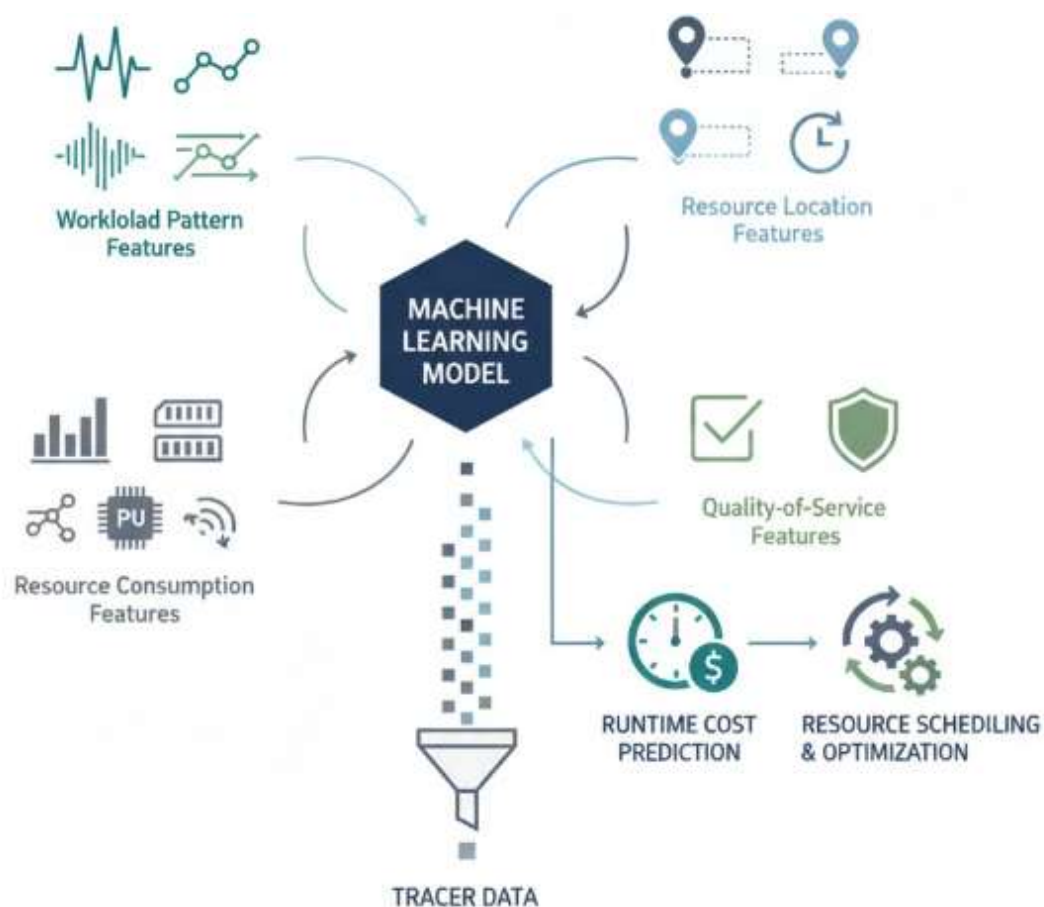


Fig 3: A Multi-Dimensional Feature Engineering Framework for Holistic Resource Allocation and QoS-Aware Scheduling in Distributed Data-Flow Engines

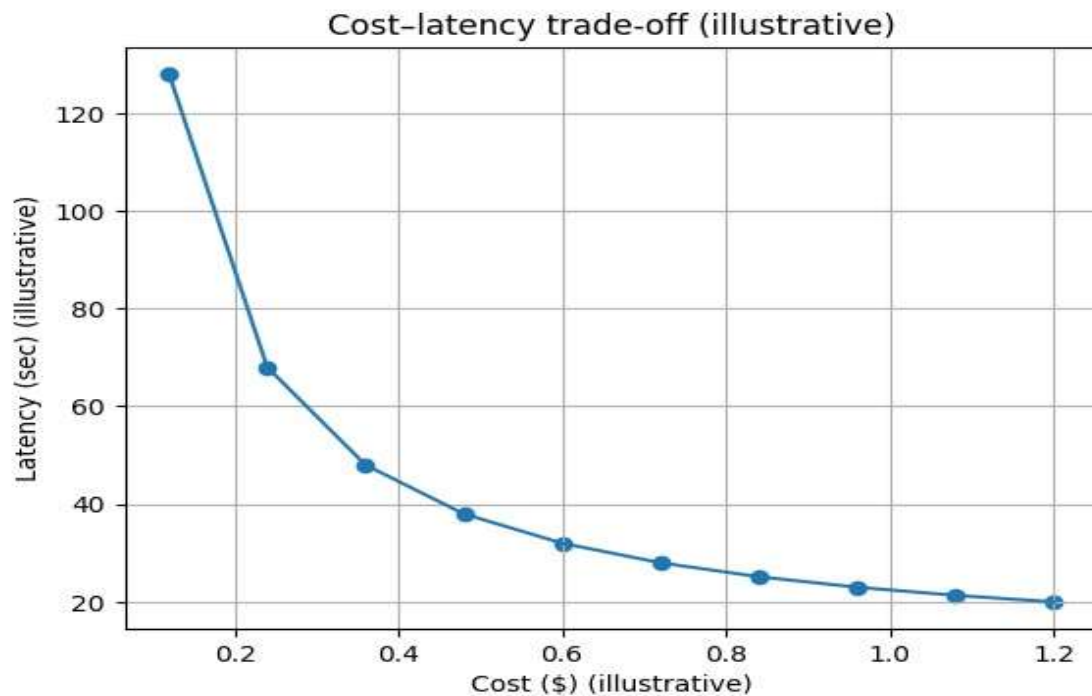
5.2. Models for Resource Allocation and Scheduling

A range of ML approaches, notably in computer vision, reinforcement learning, and time-series forecasting, present themselves as candidates for optimizing pipeline performance. Two categories stand out. SP algorithms presume the optimal solution for all distinct combinations of



input constraints resides in an external repository. For active RL and Markov decision processes, the decision-making intelligence can be internal to the data pipeline; learning proceeds via repeated interaction with the environment. The training is accomplished offline when a high-fidelity multicloud simulator of the proposed framework is available. Supervised predictors appear particularly suited for task scheduling and function placement in a data pipeline because of the multiprocessing architecture and large volumes of training data that will be generated. Training relies on workload characterization enriched with QoS requirements and resource utilization patterns and levels, complemented by resource budgets, deployment and scheduling latencies, cloud-provider specifications, resource allocation policies, and service-level agreements (SLAs). Health monitoring of the service mesh coupled with system-layer observability data will yield training and validation labels. The chosen modeling approach will thus depend on the level of feature engineering and available training and validation data.

ML also supports dynamic and adaptive optimization, balancing performance as effectively as possible while operating under continuous or transient changes in conditions; it does so by steering the scheduling and resource allocation functions toward one or more of the following goals: lowering execution latency (or latency for a batch of requests); maximizing throughput over time within an SLA (considering resource budgets); maximizing the number of servable request paths over time (or servable paths at a given instant) within SLAs and resource budgets; minimizing the end-to-end cost of the pipeline; minimizing the execution cost of the pipeline (equivalent to cost per request when operating at maximum throughput). Data-driven and AI-based techniques (including system synthesis and ML) address scheduling such that performance is improved or path latency reduced without triggering alert conditions.



6. System Design and Implementation

The supervised prediction model utilizes randomly generated workloads covering the decision space as training data, while a reinforcement learning approach learns resource-logical patterns through interaction with the pipeline. Supervised prediction models govern batch operations and the assessment of data quality rules, and predict the number of connectors required for each phase; other batch functions must also be weighted by the loading–unloading data time. Resource localization is labeled by the administrator, and the pipeline’s control plane enriches its description with exam plans for ML embeds. The reinforcement learning agent controls the parallelism level for thermal memory-free in-memory processes, and also governs the local approach to data stash and un-stash components located on nearby nodes. The deployment phase is designed to be as simple as possible.

Although the system presented in this study is limited to the scheduling of data pipelines, all processes of the described architecture and the declared end-to-end process can benefit from auto-tuning of input parameters. Models explore the usage of data locality, the resource



trajectory, and the time to finish the process against the operations being executed at all process locations through direct interactions. Such models are particularly suited for reinforcement learning. Patterns arising from the combination of pipeline components, detected by a model considering all system logs, can also compose a Markov decision process that tracks the use of resources on the system graph. Finally, the QoS requirements expressed in declarative schemas can be 0-1 labels that predict the local and global backpressure in order to generate a near-optimum routing and an optimum level of resource parallelism.

Equation 3) Core optimization problems (cost–latency trade-off)

3.1 Constrained optimization: minimize cost subject to latency QoS

$$\min_{\mathbf{r}} C_{\text{tot}}(\mathbf{r}) \quad \text{s.t.} \quad L_{\text{e2e}}(\mathbf{r}) \leq L_{\text{max}}, \quad r_i \in \mathbb{Z}_{\geq 1}$$

3.2 Lagrangian (step-by-step)

Step 1 (write Lagrangian):

$$\mathcal{L}(\mathbf{r}, \eta) = \sum_{i=1}^N p_i r_i T + \eta \left(\sum_{i=1}^N \left(\frac{1}{r_i \mu_i - \lambda} + d_i \right) - L_{\text{max}} \right)$$

where $\eta \geq 0$ is the Lagrange multiplier.

Step 2 (differentiate w.r.t. each r_i , continuous relaxation):

$$\frac{\partial \mathcal{L}}{\partial r_i} = p_i T + \eta \cdot \frac{\partial}{\partial r_i} \left(\frac{1}{r_i \mu_i - \lambda} \right)$$

Now,

$$\frac{\partial}{\partial r_i} \left(\frac{1}{r_i \mu_i - \lambda} \right) = - \frac{\mu_i}{(r_i \mu_i - \lambda)^2}$$

So:

$$\frac{\partial \mathcal{L}}{\partial r_i} = p_i T - \eta \frac{\mu_i}{(r_i \mu_i - \lambda)^2}$$



Step 3 (set to zero at optimum interior point):

$$p_i T = \eta \frac{\mu_i}{(r_i \mu_i - \lambda)^2}$$

Step 4 (solve for r_i):

$$(r_i \mu_i - \lambda)^2 = \eta \frac{\mu_i}{p_i T} r_i \mu_i - \lambda = \sqrt{\eta \frac{\mu_i}{p_i T}} r_i = \frac{\lambda}{\mu_i} + \frac{1}{\mu_i} \sqrt{\eta \frac{\mu_i}{p_i T}}$$

6.1. Data Ingestion and Processing Layer

Pipelines manifest in either streaming or batch paths. Streaming flows transfer data continuously, incurring negligible delay before reaching consumers. Streaming ingestors—e.g. Apache Kafka Connect and Amazon Kinesis Data Firehose—serve data as soon as it’s available. Sources must apply backpressure, preventing data production from exceeding transfer capacity. Latency budget dictates whether shortcuts may be taken at the expense of strong consistency; for replica pools, strong read-after-write consistency requires all writes to read from a single writer, impacting availability. The throughput-per-query bottleneck for a service mesh executing data enrichment may govern parallelism setup, promoted—or demoted—relative to throughput provided more direct paths connecting sources to consumers.

Batch jobs read data from file systems; data-in products like Amazon S3, backed by CDNs or caches, offer data-serving latency at the same order of magnitude as data-transfer latency. Natural groupings of clients with similar locality requirements allow for a drop in streaming parallelism in favour of batched requests. The batch-oriented nature of data lakes allows temporary data-discrepancy models to linearly reduce storage costs and serve data-read requests with different levels of data quality; rules of thumb for creating dedicated datasets in a layered storage architecture, plus a set of cost-performance benchmarks, improve cost optimization exponentially. A cloud-agnostic streaming engine like Apache Flink, supporting both streaming and batch modes, would facilitate the definition of autoscaling policy rules that account for backpressure, support the dynamic addition of new data sinks, and enable datacentre-best-positioned operations—e.g. ETL/ELT data-moulding operations requiring interaction with data consumers—to run close to clients-minimizing costs on data access.

6.2. Control Plane and Orchestration



The control plane is responsible for declarative descriptions of the data pipeline, such as schemas, data flow policies, and ML model metadata. Policy engines analyze these abstractions and trigger adaptation and resource provisioning operations. A dedicated set of rules evaluates resource, connectivity, and QoS requirements triggered by the pipeline definition; the results feed the autoscaling unit for proper deployment of the connectors and transformations. Data quality rules, such as completeness or timeliness, are instantiated in the data storage system, often as SQL-based checks in the defined monitoring environment. The policies implementing over- or under-provisioning of resources rely on a reaction matrix based on workload patterns detected over time. Failure and incident events caught by the observability layer drive archival procedures. Finally, the CI/CD scheme for ML components enables robustness of the data similarity-based data-locality and scheduling predictors.

For produced data in a streaming manner, a backpressure management module controls the source speed while preventing data loss in sinks. The streaming parallelism strategy relies on a user-defined factor for the connector and a per-node inbound throughput estimation multiplied by the pipelines at each level. The ability to dynamically choose the processing engine is a crucial feature of the architecture, allowing performance tuning and cost reduction. In the polling procedure, both the latency and the expected execution consumption determine whether the batch engine is engaged. For installations relying on cloud-service connectors, the selection depends on the connector's price. Finally, the choice of streaming engines follows three criteria: the ability to handle exactly-once processing, the supported processing patterns and functions, and the integration support for the CICD platform.

7. Conclusion

Data pipelines within cloud deployments can automatically adapt to workload evolution and resource availability by employing ML on relevant historical artifacts. The end-to-end architecture supporting this concept is modular, ensuring compatibility with all major cloud providers. Key components are either observability extensions to existing cloud services or implementations of well-established design patterns. Automatic behavior tuning relies on reinforcement learning, while predictive ML techniques support high-level scheduling decisions. Empirical evidence from a representative cloud provider points to an effective realization of these optimizations.

The objectives stated above are met, and important practical steps for real-world deployment and integration of the proposed ideas into a production-grade system are explored.



Results from a streaming ETL application reveal a limited degree of optimization in a single cloud environment and suggest that a unified solution across heterogeneous clouds remains an unmapped territory. Gaps in the experimental evidence thus lead to a clear set of directions for future work. Generalization and validation of the mechanisms in an unsupervised manner are necessary, as is the inclusion of user relations in the learning processes. Specialized strategies to address data residency and privacy concerns along with emergent ML-based approaches for model self-optimization represent promising avenues for the adaptive optimization of data pipelines.

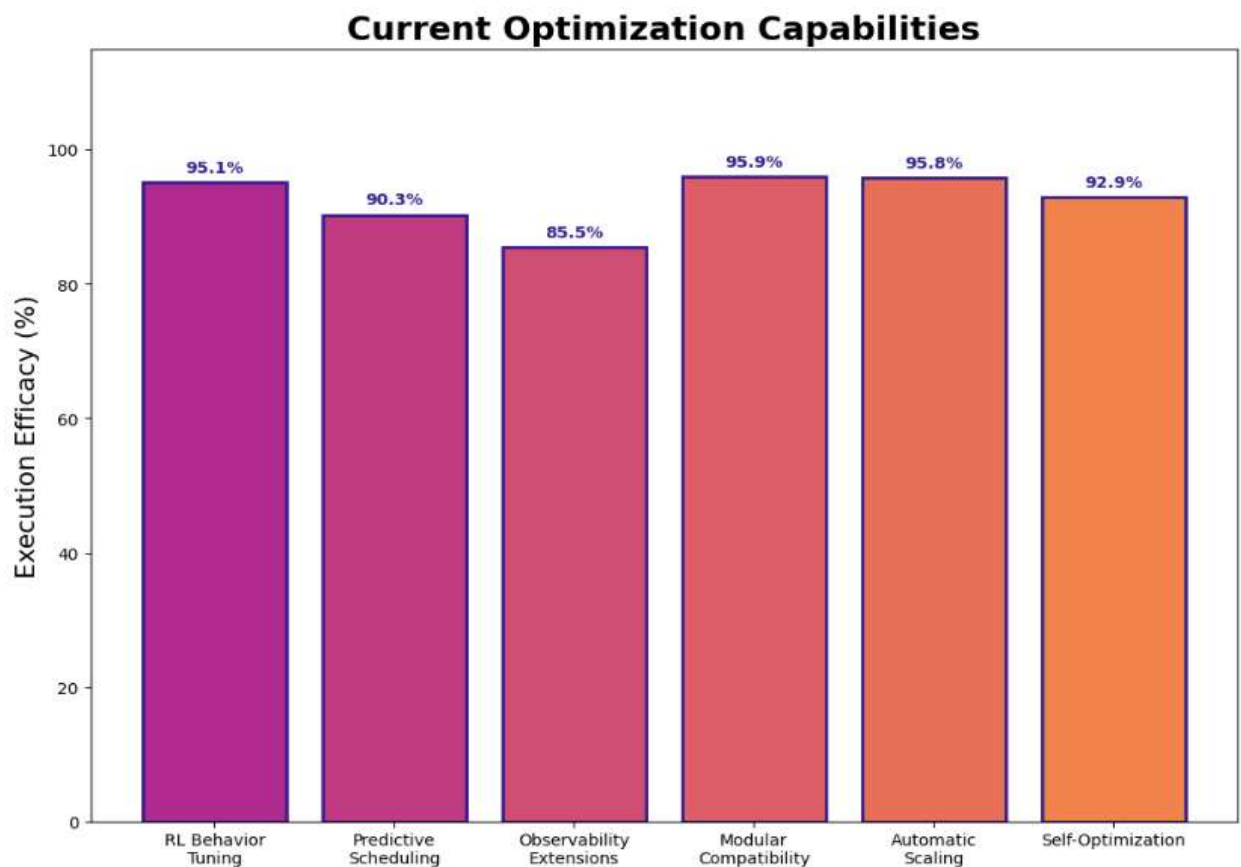


Fig 4: Current Optimization Capabilities

7.1. Future Work and Research Directions



Model generalization and transfer, cross-cloud portability, privacy protection, and the application of ongoing advancements in ML to support self-optimizing strategies represent particularly colourful yet less-explored directions. While CustardNet achieves state-of-the-art performance for the Azure data pipeline in terms of training time, pipeline cost, QoS latency, and QoS throughput, the audiovisual dataset is urban and, thus, its workload characteristics manifest in forested geographical areas as seen by the Kaliningrad cluster. Hence, CustardNet is trained solely in Azure, a cloud provider environment having non-point data locality and its visual workload. Using it on a diurnal distributed workload for urban stealthy detection on the German cluster is merely hoped to yield acceptable cost and QoS predictability without retraining, which has not yet been possible. Yet, CustardTransfer, a transfer-learning-based approach for cross-cloud data pipeline workload optimization, and the Poor-Transfer scenario—where the model pretrained on the source cloud provider environment achieves worse performance than another model pretrained on the target cloud provider environment and is nonetheless applied to mitigate the absence of training on the target cloud provider environment—seem particularly encouraging.

Experience with CustardNet in other cloud provider environments highlights risks in data privacy, especially for image classification. Reinforcement learning explores non-gaming environments. So-called emergent behaviours have developed from players' interactions with learning agents in DotA 2 and StarCraft II. Such emergent behaviours have been simulated during training but have never been achieved naturally like in CustardNet and CustardTransfer. Applying reinforcement-learning game-changing strategies to training data pipeline workload optimizers, along the lines of CustardNet and CustardTransfer, could lead to supporting players via constant automatic training.

References

1. Bandi, V. D. V. K. (2025). Self-optimizing data pipelines using machine learning for cloud workloads. *Journal of Information Systems Engineering and Management*, 10(63s).
2. Jakubik, J., Vössing, M., Köhl, N., Walk, J., & Satzger, G. (2024). Data-centric artificial intelligence. *Business & Information Systems Engineering*, 66, 507–515.
3. Kang, D., Guibas, J., Bailis, P., Hashimoto, T., Sun, Y., & Zaharia, M. (2024). Data management for ML-based analytics and beyond. *ACM/IMS Journal of Data Science*, 1(1), Article 4.



4. Davuluri, P. S. L. (2023). AI-Augmented Sanctions Screening: Enhancing Accuracy and Latency in Real Time Compliance Systems. AI-Augmented Sanctions Screening: Enhancing Accuracy and Latency in Real Time Compliance Systems (December 15, 2023).
5. Raatikainen, M., Souris, H., Remes, J., & Stirbu, V. (2024). Machine learning lineage for trustworthy machine learning systems: Information framework for MLOps pipelines. *IEEE Software*, 42(1), 51–58.
6. Zhao, M. Y., Adamiak, E., & Kozyrakis, C. (2024). Cedar: Optimized and unified machine learning input data pipelines. *Proceedings of the VLDB Endowment*, 18(2), 488–502.
7. Rad, Z. S., & Ghobaei-Arani, M. (2024). Data pipeline approaches in serverless computing: A taxonomy, review, and research trends. *Journal of Big Data*, 11, Article 82.
8. Sriram, H. K., Challa, K., Gadi, A. L., & Singreddy, S. (2025). AI and Cloud-Driven Transformation in Finance, Insurance, and the Automotive Ecosystem: A Multi-Sectoral Framework for Credit Risk, Mobility Services, and Consumer Protection. Anil Lokesh and singreddy, Sneha, AI and Cloud-Driven Transformation in Finance, Insurance, and the Automotive Ecosystem: A Multi-Sectoral Framework for Credit Risk, Mobility Services, and Consumer Protection (March 15, 2025).
9. Gu, Y., Cheng, F., Yang, L., Xu, J., Chen, X., Cheng, L., et al. (2024). Cost-aware cloud workflow scheduling using DRL and simulated annealing. *Digital Communications and Networks*, 10(6), 1590–1599.
10. Seenu, A., Sheelam, G. K., Motamary, S., Meda, R., Koppolu, H. K. R., & Inala, R. (2025). AI-Driven Innovations in Infrastructure Management with 6G Technology. In 2025 2nd International Conference on Computing and Data Science (ICCDs) (pp. 1–6). IEEE. 2025 2nd International Conference on Computing and Data Science (ICCDs). <https://doi.org/10.1109/iccds64403.2025.11209649>
11. Pan, J., & Wei, Y. (2024). A deep reinforcement learning-based scheduling framework for real-time workflows in the cloud environment. *Expert Systems with Applications*, 249, 124845.
12. Schelter, S., Guha, S., & Grafberger, S. (2024). Automated provenance-based screening of ML data preparation pipelines. *Datenbank-Spektrum*, 24, 187–196.
13. Dutta, P., Adusupalli, B., Koppolu, H. K. R., Doddla, A., Yağanoğlu, M., Banerjee, J. S., & Chakraborty, A. (2025, March). Textual Social Data Disinformation Analysis Using a Hybrid Context-Enhanced Deep Learning Model. In *Doctoral Symposium on Human Centered Computing* (pp. 342–352). Singapore: Springer Nature Singapore.
14. Majeed, A., & Hwang, S. O. (2024). Feature stores: A key enabler for feature reusability and availability across machine learning pipelines. *Computer*, 57(1), 69–74.



15. Li, G., Zhou, X., & Zhao, X. (2024). LLM for data management. *Proceedings of the VLDB Endowment*, 17(12), 4213–4216.
16. Priestley, M., O'Donnell, F., & Simperl, E. (2023). A survey of data quality requirements that matter in ML development pipelines. *Journal of Data and Information Quality*, 15(2), Article 11.
17. Narasareddy Annapareddy, V., Pamisetty, A., Malempati, M., Kaulwar, P. K., & Bhardwaj Komaragiri, V. (2025, April). Enhancing Solar Power System Efficiency Through AI-Driven Predictive Maintenance and Cloud-Based Infrastructure Stability Solutions. In *International Conference on Smart Computing and Informatics* (pp. 328-337). Cham: Springer Nature Switzerland.
18. Liu, R., Park, K., Psallidas, F., Zhu, X., Mo, J., Sen, R., Interlandi, M., Karanasos, K., Tian, Y., & Camacho-Rodríguez, J. (2023). Optimizing data pipelines for machine learning in feature stores. *Proceedings of the VLDB Endowment*, 16(13), 4230–4239.
19. Kreuzberger, D., Köhl, N., & Hirschl, S. (2023). Machine learning operations (MLOps): Overview, definition, and architecture. *IEEE Access*, 11, 31866–31879.
20. Vadisetty, R., Nuka, S. T., Kalisetty, S., Pandugula, C., Burugulla, J. K. R., & Annapareddy, V. N. (2025, February). Generative AI for Advanced Recycling Processes in Polyethylene and Polypropylene Manufacturing. In *International Ethical Hacking Conference* (pp. 269-284). Singapore: Springer Nature Singapore.
21. Faubel, L., Schmid, K., & Eichelberger, H. (2023). MLOps challenges in Industry 4.0. *SN Computer Science*, 4, Article 828.
22. Xiu, X., Li, J., Long, Y., & Wu, W. (2023). MRLCC: An adaptive cloud task scheduling method based on meta reinforcement learning. *Journal of Cloud Computing*, 12, Article 75.
23. Challa, S. R., Burugulla, J. K. R., Pamisetty, A., Challa, K., & Paleti, S. (2025, April). AI and ML-Powered Cybersecurity Strategies for Cloud Computing: Ensuring Infrastructure Stability in Financial and Retail Sectors. In *International Conference on Smart Computing and Informatics* (pp. 315-327). Cham: Springer Nature Switzerland.
24. Zhang, S., Zhao, Z., Liu, C., & Qin, S. (2023). Data-intensive workflow scheduling strategy based on deep reinforcement learning in multi-clouds. *Journal of Cloud Computing*, 12, Article 125.
25. Zhang, J., Cheng, L., Liu, C., Zhao, Z., & Mao, Y. (2023). Cost-aware scheduling systems for real-time workflows in cloud: An approach based on genetic algorithm and deep reinforcement learning. *Expert Systems with Applications*, 234, 120972.
26. Challa, S. R., Kaulwar, P. K., Koppolu, H. K. R., Adusupalli, B., & Suura, S. R. (2025, April). Big Data and AI in Wealth Management: Leveraging Cloud Computing for Secure



- and Scalable Financial Infrastructure. In International Conference on Smart Computing and Informatics (pp. 307-318). Cham: Springer Nature Switzerland.
27. Mampage, A., Karunasekera, S., & Buyya, R. (2023). Deep reinforcement learning for application scheduling in resource-constrained, multi-tenant serverless computing environments. *Future Generation Computer Systems*, 143, 277–292.
 28. Singireddy, S., Pandiri, L., Motamary, S., Kummari, D. N., Somu, B., & Lakkarasu, P. (2025, August). Blockchain-Powered Claims Validation for Enhancing Trust in Health and Auto Insurance Ecosystems. In *International Conference on Artificial Intelligence: Theory and Applications* (pp. 26-38). Cham: Springer Nature Switzerland.
 29. Mangalampalli, S., Karri, G. R., Mohanty, S. N., Ali, S., Khan, M. I., Abduvalieva, D., Awwad, F. A., & Ismail, E. A. A. (2023). Fault tolerant trust based task scheduler using Harris Hawks optimization and deep reinforcement learning in multi cloud environment. *Scientific Reports*, 13, Article 19179.
 30. Cui, T., Yang, R., Fang, C., & Yu, S. (2023). Deep reinforcement learning-based resource allocation for content distribution in IoT-edge-cloud computing environments. *Symmetry*, 15(1), 217.
 31. KOTHAPALLI, S. L., SEENU, A., DILEEP, V., & YASMEEN, Z. (2025). THE FUTURE OF CUSTOMER ENGAGEMENT IN RETAIL BANKING: EXPLORING THE POTENTIAL OF AUGMENTED REALITY AND IMMERSIVE TECHNOLOGIES. *INTERNATIONAL JOURNAL*, 73(1), 72-79.
 32. Barbudo, R., Ventura, S., & Romero, J. R. (2023). Eight years of AutoML: Categorisation, review and trends. *Knowledge and Information Systems*, 65, 5097–5149.
 33. Neutatz, F., Lindauer, M., & Abedjan, Z. (2024). AutoML in heavily constrained applications. *The VLDB Journal*, 33, 957–979.
 34. Jakubik, J., Vössing, M., Köhl, N., Walk, J., & Satzger, G. (2023). Data-centric artificial intelligence. *Business & Information Systems Engineering*, 65, 507–515.
 35. Singireddy, J., Kaulwar, P. K., Somu, B., Meda, R., Dodda, A., & Yellanki, S. K. (2025, August). Reinforcement Learning-Based Asset Allocation in Algorithmic Trading for Banking Institutions. In *International Conference on Artificial Intelligence: Theory and Applications* (pp. 357-369). Cham: Springer Nature Switzerland.
 36. Zha, D., Bhat, Z. P., Lai, K.-H., Yang, F., & Hu, X. (2023). Data-centric AI: Perspectives and challenges. *ACM Computing Surveys*, 56(1), Article 17.
 37. Mohr, F., Wever, M., Tornede, A., & Hüllermeier, E. (2021). Predicting machine learning pipeline runtimes in the context of automated machine learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 43(9), 3055–3066.



38. Ande, R., Mehta, D., Pandugula, C., Krishna AzithTejaGanti, V., & Kalisetty, S. (2025). Modeling the Som Kamla Amba sub-watershed using Artificial Neural Networks (ANN) and the SWAT tool. Available at SSRN 5341755.
39. Marinescu, R., Kishimoto, A., Ram, P., Rawat, A., Wistuba, M., Palmes, P. P., & Botea, A. (2021). Searching for machine learning pipelines using a context-free grammar. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(10), 8902–8911.
40. Orr, L., Sanyal, A., Ling, X., Goel, K., & Leszczynski, M. (2021). Managing ML pipelines: Feature stores and the coming wave of embedding ecosystems. *Proceedings of the VLDB Endowment*, 14(12), 3054–3061.
41. Maguluri, K. K. (2025). Ethical challenges in artificial intelligence. *How Artificial Intelligence is Transforming Healthcare IT: Applications in Diagnostics, Treatment Planning, and Patient Monitoring*, 132.
42. Renggli, C., Rimanic, L., Gurel, N. M., Karlas, B., Wu, W., & Zhang, C. (2021). A data quality-driven view of MLOps. *IEEE Data Engineering Bulletin*, 44(1), 11–23.
43. Patel, H., Ishikawa, F., Berti-Equille, L., Gupta, N., Mehta, S., Masuda, S., Mujumdar, S., Afzal, S., Bedathur, S., & Nishi, Y. (2021). Data quality assessment for machine learning. *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, 3141–3150.
44. Gong, F. (2021). Workflow scheduling based on mobile cloud computing machine learning. *Wireless Communications and Mobile Computing*, 2021, Article 9923326.
45. Karat, C., & Senthilkumar, R. (2022). Optimal resource allocation with deep reinforcement learning and greedy adaptive firefly algorithm in cloud computing. *Concurrency and Computation: Practice and Experience*, 34(4), e6657.
46. Gong, F., & colleagues. (2022). Reinforcement learning-assisted autoscaling mechanisms for serverless computing platforms. *Simulation Modelling Practice and Theory*, 116, 102461.
47. Zhou, R., Pang, J., Zhang, Q., Wu, C., et al. (2022). Online scheduling algorithm for heterogeneous distributed machine learning jobs. *IEEE Transactions on Cloud Computing*.
48. ADUSUPALLI, B. (2025). Redefining Financial Risk Strategies: The Integration of Smart Automation, Secure Access Systems, and Predictive Intelligence in Insurance, Lending, and Asset Management. *JAIBDD*.
49. Öztürk, E., Ferreira, F., Jomaa, H. S., Schmidt-Thieme, L., Grabocka, J., & Hutter, F. (2022). Zero-shot AutoML with pretrained models. *Proceedings of the 39th International Conference on Machine Learning*, 162, 17138–17155.
50. Mohr, F., Wever, M., & Hüllermeier, E. (2023). Naive automated machine learning. *Machine Learning*, 112, 1131–1170.



51. Pandugula, C., Ganti, V. K. A. T., Kalisetty, S., Ande, R., & Mehta, D. (2025). Modeling the Som Kamla Amba Sub-Watershed Using Artificial Neural Networks (ANN) and the SWAT Tool. Available at SSRN 5341614.
52. Subramanya, R., et al. (2022). A systematic review of machine learning operations and MLOps practices. Proceedings of relevant international software engineering and machine learning research venues.
53. Pölöskei, I. (2021). MLOps approach in the cloud-native data pipeline design. *Acta Technica Jaurinensis*, 14(4), 385–395.
54. Barrak, A., Petrillo, F., & Jaafar, F. (2022). Serverless on machine learning: A systematic mapping study. *IEEE Access*, 10, 99337–99352.
55. Faubel, L., Schmid, K., & Eichelberger, H. (2022). A reference architecture for the operationalization of machine learning models in manufacturing. *Procedia CIRP*, 115, 130–135.
56. Auer, A., et al. (2022). Qualitative assessment of the impact of manufacturing-specific influences on machine learning operations. *Procedia CIRP*, 115, 136–141.
57. Warg, M., Brazhenko, D., Somenkova, A., Golovkov, A., & Sergunin, I. (2025). MLOps architecture as a future of machine learning. *Journal of Computer Information Systems*.
58. Nogare, D., Silveira, I. F., Banzai, R., & Alexandre, M. C. (2025). Make or buy strategy for machine learning operations—MLOps. *Anais da Academia Brasileira de Ciências*, 97(2), e20240924.
59. Wei, H., García Pañeda, X., & Salvachúa Rodríguez, J. (2025). Optimizing machine learning operations in multi-cloud infrastructure: A framework for unified deployment management and topology discovery. *Cluster Computing*, 28, Article 933.
60. Kramer, K. M., Restat, V., Strasser, S., Störl, U., & Klettke, M. (2025). Towards next generation data engineering pipelines. Proceedings of the VLDB-related data engineering research community.