

Chapter 2: Architectural Principles for Cloud-Scale Data Systems

2.1. Introduction

Cloud-scale services exhibit an increasing diversity of workloads and user demands, reflected in a growing range of deployment models by the service consumers. Multitenancy, autoscaling, performance isolation, energy and carbon footprint reduction, and lifecycle-management support are now a requirement rather than a luxury. Architectural trends clearly point in the direction of service-oriented or microservices architectures. However, casuistics about cloud-scale services also evidence a wide range of rarely discussed data-centric aspects of their architectures. Addressing these considerations in a structured manner can help cloud-scale data systems evolve toward appropriate common patterns and paradigms that ease their adoption and operation and constitute sound foundations for training cloud computing practitioners.

Four groups of cloud-scale deep architectural considerations can be defined: design aspects that are manifest at the system level and affect its usability and operation, data-centric aspects of the architecture supporting the service function and fulfilling cloud-scale requirements, the data management subsystem responsible for data lifecycle and quality, and the storage-and-retrieval and compute-and-processing foundations of the system. Main design aspects involve provision of dedicated natural-resource pools at a reasonable cost and minimal energy/carbon footprint, a scheduling mechanism that attains workload isolation while efficiently deploying shared services, and support for Green computing and the triple bottom line. Supported cloud-scale data workloads cover independent multisource data sharing on a large scale, data pipelines that combine analytics and operational functions in a unified flow, and long-term storage of huge amounts of rarely accessed data.

2.1.1. Capacity Planning and autoscaling

The cloud model offers a fog of virtual resources to users and applications. Resources are allocated and released in response to demand. But the fog does not match precisely the demand all of the time. Consequently, managed resources must be provisioned with particular attention to costs. Here, fitting predictions of demand and supply are required. Because resources are supplied by third parties, it is not always possible to get needed resources at all times and therefore the elastic model becomes very important for systems of cloud-scale. Resources at runtime must be allocated to specific workloads in response to changing demand but the scales must remain independent, implementing workload isolation. Such resource scheduling can efficiently be improved using free and preemptable resources such as those offered by Spot. Also important for cloud providers is lifecycle management. Managed resources have their own lifecycle that must be monitored, starting from provisioning new ones, keeping them available when used, stopping unused for long, decommissioning broken ones, those with increasing costs, and at the end removing ones never used or that barely are not enough cost-effective.

The electricity needed to power clouds can be derived from residential power sources, but besides the large amounts of needed sources electricity can become an issue for clouds too. Providing a cloud service that is green and environmentally friendly is becoming more and more a business driver for hosting clouds providers. Clouds consuming more power are less green and therefore various initiatives for finding solutions to all aspects of consumption of electricity are being researched. Cost control is a big issue in servers as well: systems always running costs more than being shut-down; system maintenance and support require higher levels of human resources. Again part of the increase of costs can be minimized or completely solved by accepting stop services making use of free and/or internal-used resources.

Equation 1: Scalability Equation (Throughput vs Resources)

$$\begin{aligned} Throughput &\propto Nodes \times Efficiency \\ Throughput &\propto Nodes \times Efficiency \end{aligned}$$

Interpretation:

- Increasing the number of nodes (horizontal scaling) increases throughput
- But real systems are limited by coordination overhead, network latency, and inefficiencies

2.1.2. Workload Isolation and Resource Scheduling

In cloud computing, workloads belonging to multiple users, teams, departments, or projects often share resource pools. This multi-tenancy approach transfers infrastructure management overhead from users to the cloud provider, minimizing dedicated resource allocations. However, multi-tenancy from a single user perspective can lead to disparate priorities and compromise quality-of-service objectives; hence, isolating workloads with dedicated resources can be beneficial. Workload isolation and resource scheduling are therefore vital cloud-scale infrastructure architecture features and typically managed as part of the cloud orchestration layer.

Workload isolation can take two forms, resource reservation and resource dedication. Typically, virtual machines are reserved for specific tasks, ensuring that a predetermined capacity is available when required, while clusters in serverless compute environments are dedicated to specific workloads. Users concerned about resource over-commitment and contention latencies can configure reservation levels. Network I/O can also be isolated using virtual network functions. Resource isolation delivers predictable performance, but resource-utilization levels decrease, generating potential wastage. Resource scheduling reduces such wastage by allocating resources to a specific workload only during the active period. Resource scheduling is typically performed in the orchestration and virtualization layers.

2.1.3. Green Computing and Lifecycle Management

Green computing emphasizes reducing power consumption at scale and designing a system that balances cost with environmental impact. The manner in which applications and users function also influences this aspect. Scheduling and lifecycle management of resources, such as autoscaling, are crucial, as is selecting the least expensive and most environmentally friendly location for processing workloads and storing data. Simple point solutions exist, such as opting for a public cloud service when deploying an application. Survival of the fittest principle still holds. It remains essential to monitor the cost incurred by a workload and choose a location that minimizes it. The cost can be associated with: (1) processing and storing data, (2) location, (3) the least environmentally damaging provider, and (4) compliant cloud providers.

Lifecycles extend beyond applications and data to the entire stack down to the platform level. Cloud service providers assume responsibility for the infrastructure and even large portions of the platform. A service may be successful and attract unexpectedly high volumes of usage, making its footprint disproportionately large. Users should be encouraged to migrate from early deployments to scale-optimized versions. For example, a good log store might employ a cloud facility optimally configured for

logging; however, when the volume decreases, it is silly to maintain that cloud facility and services. Smaller users of a cloud service might prefer data going through the main service, while larger users with specific characteristics might use another cloud provider. Eventually, the specialized providers may become irrelevant, yet their services should be maintained during their popularity, while other specialized services appear for a while.

2.2. Architectural Paradigms for Cloud-Scale Data

Many cloud-native and scalable data systems are partially or fully distributed, rely on data-oriented paradigms for ramping up their maintainability, and provide support for various processing capabilities—batch and stream, real-time and deferred—either as part of their core design or as add-ons. These properties influence the provisioning phases of cloud-scale data systems and can guide the selection of specific design approaches that address particular scaling and operational concerns:



Fig 2.1: Architectural principles for cloud systems

Most data systems and applications are hosted as collections of services running in virtual or containerized environments, facilitating work isolation, languages and framework diversity, and independent lifecycle management (continuous delivery). Service-oriented architecture (SOA) has become a pervasive approach for building

cloud-native applications, with microservices being a more focused manifestation addressing independent deployment, development, and operation.

A Snowflake case study illustrates autoscaling and the provisioning of processing resources as a function of workload. Data management support at scale encompasses schema evolution, storing and querying heterogeneous data sets, deriving quality and lineage metadata, and other processes.

2.2.1. Service-Oriented and Microservices Architectures

Automation and elasticity allow cloud-scale data systems to increase resource capacity beyond the physical limits of a single data center. Such clouds serve arbitrary workloads at a global scale, typically designed to ensure high availability rather than peak performance on a specific task. Consequently, bottlenecks must be avoided in software design and management. Service-oriented architecture allows a workload to be divided into a set of independent services, instantiated and scaled independently. Other macroservices can cater for batch workload by implementing the lambda architecture pattern where replayable events are stored in fast, immutable, cheap-to-store sequences and periodically aggregated for serving. Microservices architecture is a natural extension, reminiscent of Kent Beck's Extreme Programming concepts applied to coding, where every service can choose its development languages and frameworks independently, thus allowing continuous delivery and aligning with the deployment economy-of-scale principle.

Designing cloud-scale data systems using a data-centers-centric view leads to two key economic principles. The first principle states that any workload requiring more infrastructure resources than a single data center can provide for the highest demand at a reasonable cost should preferably be deployed across multiple data centers. The second principle states that the higher the global bet capacity of the computation resource, the higher the global bet capacity of the entire resource group. Both principles should be considered when partitioning data. A partitioning strategy similar to Kafka's topic-partitioning should be adopted in order to achieve seamless dynamic scale-up and scale-down of read loops, which are mainly impacted by the number of partitions.

2.2.2. Data-Centric Architectures and Data Mesh Concepts

Data-centric architectures, which advocate the design of applications around data rather than around an application framework, can improve development efficiency and quality. Communication across boundaries is typically minimized or segregated, thus fostering team ownership and facilitating deploying and scaling in isolation. Conventional

database federation, however, fails to adequately serve heterogeneous needs. In broadening the domain of polyglot persistence, the concept of data mesh formalizes domain-oriented, decentralized data ownership and architecture in a self-serve product line for full-cycle data products.

Data mesh broadens into a nondogmatic application of a principled data product governance framework and good data product management, proportionate to polyglot persistence domain decomposition. Distilling the conventional perspective on polyglot persistence data systems yields paralleled classifications into dedicated slices and into distinct products, informed prudently by the combined perspectives of traditional outsourcing, data product companies like Google and Pinterest, and general data product principles. The bifurcation naturally generalises into a data mesh offering a data product product line building on outsourcing principles.

2.2.3. Lambda and Kappa Patterns for Stream and Batch Processing

Event-driven architectures that respond to incoming data in real-time offer an alternative to complaint-driven systems that undertake batch actions at predetermined intervals. Popularized within the realm of cloud-hosted systems processing data at scale, the lambda pattern exploits highly scalable and elastic stream-processing frameworks alongside fault-tolerant batch-processing components. Dedicated microservices present on dedicated infrastructure automate pipeline components, ensuring appropriate actions on incoming data. Although originally named for AWS cloud services, analogous paradigms have emerged elsewhere (GCP, Azure, IBM, Snowflake, ...) independently of the cloud provider.

Lambda Architecture (LV) merges the merits of both batch and stream processing. Batches applied to the complete dataset (using more costly dedicated microservices) build a materialized view, while incremental materialized views propose an alternate approach. Batch and stream results are merged to produce the final outcome. Addressing the original challenges of Lambda Architecture, Kappa Architecture proposes a single-stream processing layer that maintains sufficient data persistence for replaying the complete dataset upon code modifications.

2.3. Data Management at Scale

Three principles specific to data management emerge from experience architecting cloud-scale data systems: embrace polyglot persistence to optimize storage and retrieval performance; invest in systems and processes to ensure data quality, provenance, and

lineage; and recognize that storage and retrieval are merely the foundation upon which service- and application-related development and operations must be built.

Polyglot persistence – the notion of using diverse data management technologies for different types of data and different classes of workload – is a highly effective way to ensure that cloud-scale applications meeting the requirements end-users and customers expect. The Principal Architect at one of the world's largest social media platforms observed that "no single database can effectively support all application functions," emphasizing the importance of choosing the right tool for the job. Nevertheless, the layered abstraction technology is just a single point solution. One of the major cloud vendors boasts "over a hundred" different database services, and with good reason: Whether as fully managed, or as a service through either code-based configuration management or CWCMS systems, a multi-tenant database such as a relational database is typically used for blobs of transactional data, caching puppy pictures, or a discovery index, while other kinds of databases and non-database technologies are often employed for telemetry data, fast text search, or time-series data.

Contrary to the views of that Principal Architect, the extensive variation of structural requirements, type of access patterns, latency targets, the set of consistency and durability guarantees required, costs and risk tolerance make polyglot data management almost a necessary condition for many business workloads. Cloud vendors and cloud-service providers invest heavily in metadata management, as efficient metadata management capabilities are key to the cost and speed of storage, retrieval, and analytics, and underpin advances such as Unified Data Mesh. All three main categories—search engines, data catalogs, and metadata management systems—are important elements in the efficient operation of these systems and should be considered in any cloud-scale architecture.

Equation 2: CAP Trade-off Equation

$$\begin{aligned} & \textit{Consistency} + \textit{Availability} + \textit{Partition Tolerance} \\ &= 2\textit{Consistency} + \textit{Availability} + \textit{Partition tolerance} \\ &= 2\textit{Consistency} + \textit{Availability} + \textit{Partition Tolerance} = 2 \end{aligned}$$

Interpretation:

- In distributed systems, you can only fully guarantee two out of three:
 - Consistency (C)
 - Availability (A)
 - Partition Tolerance (P)

2.3.1. Data Modeling for Polyglot Persistence

Cloud-scale data systems must uniquely handle the peculiarities resulting from applications and services that exploit a common shared set of data. Data management for cloud-scale data systems does not mean a monolithic or homogeneous approach to data storage, but rather applies the idea of polyglot persistence to provide the right data store for the right problem. This introduces complex architectural and engineering requirements with respect to data modeling, where inconsistency and data sharing have a different perspective and requirements from traditional data management.

In polyglot persistence architectures, schema design varies according to the read and write patterns from services and applications, what data needs to reside close together to minimize retrieval costs, how data is loaded into a store, and how data is queried. Inevitably this leads to redundancy, both structural and in terms of data content, but computations over individual stores can be simple and fast, resulting in a reduction in computational complexity and cost. The management and supply chain of shared data requires data discovery capabilities, defining responsibilities for maintaining data consistency, and providing visibility and monitoring of data quality. Key areas for attention in the architectural design of data modeling includes: Polyglot Persistence with Data Sharding and Replication, Metadata and Governance Information Management, and Data Quality, Provenance, and Lineage.

2.3.2. Schema Evolution and Metadata Management

Auto-generated text using the supervisor's instructions follows.

Different storage engines have different capabilities and tradeoffs, a pattern known as polyglot persistence. However, the result is greater operational complexity, and many cloud-scale data systems include support for multiple storage engines. The addition or removal of chosen database engines must be conducted with caution, especially in the Data Centralized and Data Mesh Paradigms, as different engines may have had different development lifecycles that may elevate time-worn data quality issues.

Schema evolution is a core component of lifecycle management for cloud-scale data systems, yet also taxing in terms of metadata management, as both the data and its associated metadata need to be updated. Schema evolution steers users and application developers towards forward-compatible schemas and aids automated enforcement of schema-on-write correctness. Support for schema evolution is a key criterion for data engineering frameworks and cloud data warehouses. Although serverless compute architectures promote a data-centric, unifying approach to data engineering and analytics, coupling serverless compute with a multitude of diverging data processing engines constrains the evolution of the core layer metadata.

Although primarily targeted at cloud data lakes, CloudEthos and Altinity provide glances of other uses for Unified Data Mesh. CloudEthos augments the core metadata management with an arbitrary set of expressible data regulations along data-provenance and data-lineage exploratories, while Altinity promise analysis of a particular type of system downtime.

2.3.3. Data Quality, Provenance, and Lineage

Data systems must ensure data accuracy and correctness, especially as systems evolve, through concepts such as data validation, data profiling, and data cleansing. Provenance and lineage concepts allow understanding how data was obtained and how it has changed, especially when it comes to audits and compliance.

Data quality addresses the degree of usefulness of data, which depends on its contextual appropriateness. Various dimensions exist to assess data quality, including accuracy, completeness, consistency, and temporal and spatial validity. The dimensions originating from the practical need of the users have resulted in different data quality management processes. Data validation applies specific validation rules to check whether data in a record satisfies requirements. Data profiles store curated summary statistics to help understand data quality conditions, inform data consumers, and act as an additional data quality dimension when compared against reference profiles.

Data cleansing may contain removal or filling of missing values, generalization or specialization of values, correcting erratic values, identification of clustering and outlier values, as well as identification and correction of duplicates. Provenance refers to data's origin while lineage captures the data's life cycle. Activity monitoring and recording, combined with structured and unstructured annotations, enables provenance capture for any existing data and its processing activity. Provenance and lineage are mainly applicable in the context of audit and compliance requirements.

2.4. Storage and Retrieval Layer

Research on cloud-scale data systems has produced a variety of new storage technologies for structured, semi-structured, and unstructured data. But their sheer number and level of complexity can pose significant challenges to designers. A mindset for adopting numerous specialized data management technologies has emerged in the form of polyglot persistence, which allows data to be stored and managed in the most sensible manner for both performance and maintainability, irrespective of whether the technology was developed in-house or acquired from an external source. Automatic data selection and movement across different management systems through ETL processes are highly

valuable for data science and artificial intelligence workloads. New storage components include cloud object storage, wide-column stores, time-series databases, distributed file systems, data lakes, search engines, and in-memory analytics engines for accelerating analytical workloads.



Fig 4.2: Cloud-scale data architecture overview

2.4.1. Object, Wide-Column, and Time-Series Stores

Numerous storage systems that scale horizontally are collectively termed NoSQL databases. Data models differ among these systems: key-value, document, graph, object, column-family, wide-column, and time-series. Key-value stores associate keys with binary objects or values. Documents encapsulate objects in specific formats, primarily JSON and XML. Graph stores enable traversals among nodes linked in various ways—directed, undirected, multigraph, weighted (for relationships), and more. Object stores manage unstructured data as keys with associated blobs. Time-series stores organize time-dependent entries for monitoring and other logging needs. Column-family databases, grouped for speed, partition rows into column families that reside together. Wide-column databases (e.g., Apache Cassandra) enable record-level customization by allowing tables to store dissimilar columns.

Three classes of NoSQL systems take center stage: object stores (represented by Systems like Amazon S3), wide-column databases (especially those built on Bigtable or Dynamo principles), and time-series stores. Broadly accessible via Cloud Storage APIs, key-value systems like Amazon S3, Azure Blob Storage, and Google Cloud Storage provide diverse capabilities and heterogeneous service portfolios. Fundamental responsibilities

include data retrieval, data loading, data lifecycle management, data organization into a hierarchy, caching, and serving as a syslog for resource audit trails. Transactional guarantees must be weak to cope with multi-tenancy. Integrating with other Google Cloud Platform services, Cloud Storage intelligently shards data and supports both regional and global data location across endpoints. Cloud Storage keeps data across its entire lifecycle, from creation to archival to deletion, with integrated Cloud Storage Lifecycle Management functionality.

2.4.2. Distributed File Systems and Data Lakes

Distributed file systems support large numbers of clients accessing very large files concurrently. They are built from clusters of commodity servers whose only purpose is to offer a single namespace for those files. The systems run without any single point of failure, rely on data replication and erasure coding for reliability and durability, federate storage and retrieval of very large files across multiple data centers, and support phasing hot data into memory, disk, and cloud.

A data lake serves as flexible storage for organization-wide analytics workloads. It streams ingest and amplifies the architectures for Spark, Flink, Presto, Athena, and similar services. It supports federated queries across multiple data stores. Aliyun have commercialized the Amazon data lake solution.

The data acknowledged the logical differences between structured and nonstructured data. A full schema can impede ingestion for complex data and formats; schemas stored externally or applied late can enable faster data loads. The principles also advised lateral connections between storage managers and application developers to maintain data hygiene and appropriate governance as the lake grows.

2.4.3. Indexing, Search, and Analytics Acceleration

Search engine technology is extensively employed in cloud-scale data systems, both internally and in user-facing connections, to enable fast, full-text retrieval from vast information collections. Specialized data stores exploit inverted indexing structures to deliver rapid responses; however, having dedicated stores for such requests results in duplication of data and substantial potential for inconsistency as the information transforms in distinct pipelines. Efficient search over heterogeneous data appearing within a system therefore often requires additional solutions to alleviate these issues, primarily through incremental, background population or via dynamically augmented pre-computed indices.

Another area within the indexing space relates to improving the performance of analytical workloads, where a small set of columns need to be searched over a large number of rows. Solutions here include the building of bloom filters that reside in memory and that can ascertain which shards likely do not contain the queried information. External systems such as the Apache Parquet format can replicate some of this functionality on disk by preserving per-column metadata through Min/Max value analysis for every column on a per-row-group basis. Various other more complex forms of optimization also exist specifically tailored at expediting analytical request processing against large datasets.

2.5. Compute and Processing Foundations

Evidence-based arguments, formal architecture, and clear reasoning support the conclusion that cloud-scale data systems demand distinctive properties. The architecture of cloud-scale data systems is driven primarily by nonfunctional requirements derived from large-scale systems operation, data engineering, and cloud infrastructure.

The compute and processing foundation reflects the need to decrease complexity for application developers. Serverless architectures decouple the execution of application code from infrastructure management, enabling workload auto-scaling at the function level. Stream processing architectures support the ingestion of continuous data feeds generated by user and system interaction, business activity, the Internet of Things, and machine-generated sources. Such systems holistically support complex event processing, monitoring, and alerting. The distinction between processing streams of data in-flight and batches of data-at-rest emphasizes the need for cloud-scale data systems to adopt batch and streaming patterns, orchestrating data processing with separate pipelines and systems. Batch processing is organized into a logical pipeline whose execution is orchestrated by a workflow engine, whereas orchestration is needed for the management of resource-intensive and long-running workflows.

Serverless compute is suitable for workloads requiring burst capacity and small-scale jobs with low-frequency execution. It offers cost advantages when code execution times are relatively low, frequently removing and re-creating resources. Bandwidth-limited compute domains and services requiring low-latency access to state are unnatural fits, as is large-scale data processing where the startup latencies of cloud resources are in the order of minutes. The nature of the data-adaptive control patterns for batch execute clouds lends a natural fit for dedicated compute patterns.

Equation 3: Data Latency Equation

$$\begin{aligned} \text{Latency} &= \text{Processing} + \text{Network} + \text{QueueingLatency} \\ &= \text{Processing} + \text{Network} + \text{QueueingLatency} \\ &= \text{Processing} + \text{Network} + \text{Queueing} \end{aligned}$$

Interpretation:

- Total latency is the sum of:
 - o Processing time (compute)
 - o Network delay (data movement)
 - o Queueing delay (load/congestion)

2.5.1. Serverless versus Dedicated Compute

Serverless compute is a managed service model usually offered for event-driven workloads and based on the extreme virtualization of infrastructure resources (compute and memory) into small containers. Resources are metered at a fine granularity and pause automatically when not in use. Applications are not made of long-lived servers but are comprised of individual stateless functions which interact through APIs or message queues with any external services, including data storage. Serverless function architectures, which leverage managed queuing, notification, and orchestration services, are also proposed as the ideal building blocks for multi-cloud systems. However, concerns remain about cold start performance, pricing transparency, and resource sharing at such fine granularity.

Dedicated compute is suited for workloads that demand guaranteed low latencies and predictable performance (such as interactive services), for sustained and higher utilization levels, or are not well suited for extreme resource virtualization (like large in-memory or GPU-heavy jobs). The trade-off is a higher operational overhead, typically solved through specialized scheduling over time-shared pools. Servers can be assigned in auto-scaling groups and throughout scheduling pipelines with provisioning times appropriate to the workload. Dedicated compute resources may also be reclaimed for batch workloads within a time-sharing model.

2.5.2. Stream Processing Architectures

High-throughput and low-latency data pipelines, and event-driven microservice architectures with message brokers, benefit from a dedicated infrastructure for stream processing. Stream processing engines, in contrast to messaging middleware, focus more upon guaranteed delivery, when and in which order events are processed and the

semantics of event processing such as time-window aggregation etc. Stream processing is a well-accepted architectural pattern but latency and throughput must be considered when designing data processing pipelines.

Stream processing has become more flexible with the emergence of a number of scalable engines supporting arbitrary pipelines beyond simple event routing. These parallel engines support fault tolerance, scaling up and down with the workload, continuous processing, windowing and event-time semantics; and some offer state management. Open-source stream processing frameworks such as Apache Storm, Apache Samza, Apache Spark Streaming, StreamX and Apache Flink, and commercial streaming services such as AWS Kinesis and GCP Dataflow offer glue in a microservice architecture that support data analytics, AI and machine-learning workloads. Nevertheless, the paradigms supported are still those envisioned when messaging first appeared: messaging middleware laid the foundations for the actor model of computation while stream processing was an instantiation of DataFlow. Indeed, the design of both takes into consideration the same context: parallelism on multi- and many-cores with the memory bottleneck.

2.5.3. Batch Processing and Orchestration

Despite their underlying differences, the push to offer batch processing as a service is evident in managed services such as AWS Batch and Azure Databricks, and by solutions that provide orchestration capabilities like AWS Step Functions and Microsoft Power Automate. While the term serverless implies a lack of user-managed compute infrastructure, it encompasses services that automatically allocate, take advantage, and deallocate resources based on demand. In the case of batch processing, auto-scaling clusters is a common design for dedicated compute services as well. By handling compute resource management and lifecycle, managed solutions enable users to focus on defining the tasks to execute, orchestrating workflows and integrating results, rather than provisioning VMs, making them easier to use.

Orchestration services provide visual or code-based interfaces to automate the execution of sequences of tasks or processes that can span multiple services, applications, or even public APIs across several cloud platforms. Such processes help define data pipelines and orchestrate ETL or ELT workflows when designed to integrate with cloud data warehouses or other dedicated processing solutions.

2.6. Networking, Security, and Compliance

Identity, access management (IAM) and secrets management represent critical success factors for cloud-scale data systems. IAM links users, groups, machines, and services to the actions they are permitted to perform on cloud resources. In the general case, there are many IAM entities with complex dependencies, including federated users that interact with an IAM system outside the cloud provider. Consequently, IAM management typically requires dedicated tools that help define policies, detect permissions subdomains, and enhance compliance readiness. For services and applications, each entity must primarily possess the minimum privileges required to successfully complete its designated job. Secrets—such as passwords, API tokens, and certificates—must remain confidential at rest and in transit, using their own storage tools, special access controls, and/or encryption. These secrets should not be hardcoded in applications and must be rotated and accessed dynamically through APIs.

Encryption is pervasive in cloud environments. Data is encrypted at rest in storage systems using a key management system (KMS) to control access to the keys, while encrypted data is decrypted transparently when loaded into a data processing engine. Data may also be encrypted in transit for a given connection using transport layer security. Data privacy at rest, in transit, and in use relies on the encryption of sensitive attributes; it is improved with further techniques such as tokenization and data masking. Nevertheless, the ability to audit when and how data is decrypted at each step in the data processing life cycle remains critical to compliance.



Fig 2.3: Cloud architecture security, compliance, scalability

2.6.1. Identity, Access Management, and Secrets

Cloud service providers offer a multitude of identity and access management features. These can be broadly classified according to the types of objects within a cloud system of managed services. Identity providers manage user identities, authentication, and authorization (the right access levels for each identity), typically through out-of-band authentication mechanisms such as one-time passwords (OTPs). Role-based and attribute-based access control (RBAC and ABAC) frameworks allow fine-grained permissions for all other managed objects (services, virtual machines, storage accounts, data sets, network components, etc.). Secrets engines, used to manage secrets such as passwords, API keys, and certificates, are especially useful for credential rotation, allowing secrets to be changed regularly and the latest version to be fetched at runtime by the relevant applications or services.

The strengths of the IAM system should, however, lead to the opposite effect, allowing no secrets needed to be stored at all within application deployment packages. This indicates possible unusual or misconfigured cloud architecture. In a fully deployed cloud ecosystem, service-to-service communication should never be performed using API keys stored in JSON configuration-like files. Such secrets-engine-based IAM features should also preclude the need for static storage accounts and keys used by services to connect to cloud storage. Services running on service providers—be they virtual machines, containers, serverless functions, or microservices—should leverage IAM systems supporting direct identity federations with storage services (in most cases and by default) to transparently access data store using native or highly optimized SDKs. Legitimate-enabled virtual machines accessing cloud object storage services using standard or SDK-based definitions benefit from advanced performance and security capabilities (protocol-level and on-demand encryption, fine-grained access control, etc.) with lower operational costs and higher availability.

2.6.2. Encryption, Key Management, and Privacy

Workload isolation and resource scheduling fall under system management, which can be seen as two parts: (1) managing the available resources to cater to the workload demands and (2) isolating the workloads that should not interfere with each other. Green computing aims to realize effective energy savings by powering on and off virtual machines based on demand periodically and using VMs with minimal load. The lifecycle of a cloud service comprises its start-up, steady-state, and shut-down phases. In the start-up phase, resources are allocated progressively as the demands grow, ideally without overshooting. In the steady-state, additional incoming demands are supported via autoscaling. In the shut-down phase, demand tapers off and resources can be reclaimed.

Lifecycle management deals with the various phases of a service. A service gets deployed with small resources. When a breach occurs, additional resources kick in and, when the breach has passed, resources can be reclaimed. Breaches or violations take place when load exceeds a defined level.

2.6.3. Compliance by Design and Auditing

End-user privacy is a source of legal and business risk for services with a cloud-scale data architecture. Laws such as the General Data Protection Regulation endow users with rights on their data, including the rights to see its contents and have it deleted. Extracting data from deep in a cloud-scale data architecture, possibly with obfuscation, is therefore an important requirement for any such service.

From a risk-management perspective, controlling user access to data is essential for limiting exposure to liable content. Legal liability extends beyond the owners of a cloud-service provider to its employees, farmers, and possibly customers. Granting the fewest privileges required for a role can be addressed with identity, access- and secrets-management services.

With significant data being exposed to public cloud providers, the possibility of encrypting data at rest, in transit, and (to the extent that it may be communicated) in use is common to cloud-scale architectures. Managing the relevant keys is essential. Encrypting user identifiers that do not need to be in plaintext when controlled by employees other than the ones consulting customers provides additional security, shifting the burden of offending data to the users themselves. In summary, addressing an architecture's vulnerability by design is crucial for cloud-scale data services. Auditing tools and processes monitoring the existence of protected-content indicators throughout the clouds, especially on data available to boundless actor, visualising data from user queries, and pattern matching for unwanted content all add noise to the actions of users browsing accepted data.

2.7. Conclusion

Cloud-scale data systems increasingly reflect existing principles of enterprise architecture and large-scale data management, suggesting that layered frameworks and distinct persona-dependent design and operational processes can contribute to a more holistic and elaborate understanding. Additionally, the increasing segmentation of cloud-scale data environments into service-oriented and microservices elements increases the need for explicit approaches to data governance and management. Noteworthy

additional areas of exploration are socio-technical, design and data role models and laws—such as Conway’s Law, Brooks’s Law and Strachey’s Law.

Future developments in cloud-scale data design and operations will be determined by the prevailing trends in enterprise architecture, large-scale data management and adjacent domains, such as artificial intelligence. Autonomy will increasingly occur at the component level while the integration of security and compliance controls into the design and build environment—compliance by design—will become mainstream practice.

2.7.1. Emerging Trends

Emerging trends in cloud computing can be observed at multiple levels of abstraction: (1) underlying infrastructure; (2) integrated offerings by cloud service providers; and (3) user adoption and awareness of cloud services. Recent advances in hardware technologies such as specialized processors (FPGAs and TPUs) and fast networking are opening up new opportunities for offering dedicated hardware acceleration for specific workloads. New protocols and standards (such as Quantum Key Distribution) are being explored to ensure secure and private communications. Service offerings by major cloud service providers are constantly expanding, especially to address top concerns in cloud adoption. Examples of such services include policies for restricting data transfers across regions to ensure regulatory compliance, simplifying identity and access management across multi-cloud deployments, automating the provisioning and configuration of cloud services, and providing managed services for compliance requirements. On the user side, awareness about cloud capabilities and limitations is increasing. Education and outreach at the user level are helping organizations make better decisions in selecting the right cloud services, architecting applications to leverage the benefits, and avoiding common pitfalls.

Focus areas such as reducing cloud costs, improving cloud security, and ensuring environmental sustainability of cloud operations and services remain paramount. Several users are starting to take responsibility and a proactive approach toward minimizing undesirable cloud costs for example, provisioned resources that are left active while not being used. A growing number of organizations are working to adopt Cloud Sustainability by Design practices so that the carbon impact of the given cloud services can be under control. Methods such as carbon-aware workload scheduling and data storage in the most cost-effective geographical region are gaining traction. As more organizations migrate important services to the cloud, risk mitigation becomes critical. A combination of cloud-native services and multi-cloud deployment strategies are increasingly deployed to provide better control and compliance assurance over cloud security..

References

- Gama, J. (2023). Knowledge discovery from data streams. *CRC Press*.
- Gottimukkala, V. R. R. (2020). Energy-Efficient Design Patterns for Large-Scale Banking Applications Deployed on AWS Cloud. *power*, 9(12).
- Hellerstein, J. M. (2023). Data systems: Principles and applications. *ACM*.
- Nagabhyru, K. C. (2024). Data Engineering in the Age of Large Language Models: Transforming Data Access, Curation, and Enterprise Interpretation. *Computer Fraud and Security*.
- Chen, X., et al. (2024). DataOps: Continuous integration for data pipelines. *IEEE Software*.
- Garapati, R. S., Adusupalli, B., Kaulwar, P. K., Gadi, A. L., Annapareddy, V. N., & Challa, K. (2025, December). The Evolution of Digital Payments: A Study on AI-Powered Transaction Monitoring Systems. In *2025 3rd International Conference on IoT, Communication and Automation Technology (ICICAT)* (pp. 1-8). IEEE.
- Sculley, D., et al. (2023). Hidden technical debt in ML systems. *NeurIPS*.
- Kolla, S. H., Inala, R., & Kumar, M. V. K. (2026). Secure RAG Architectures with Small Language Models for Governance-Aligned LLM Deployment in Enterprise Service Management Platforms. *International Journal of Economic Practices and Theories*, 2026, 166-179.
- Amershi, S., et al. (2023). Software engineering for machine learning. *ICSE Proceedings*.
- Davuluri, P. N. (2020). Improving Data Quality and Lineage in Regulated Financial Data Platforms. *Finance and Economics*, 1(1), 1-14.
- Doan, A., Halevy, A., & Ives, Z. (2012). Principles of data integration. Morgan Kaufmann.
- Segireddy, A. R. (2021). Containerization and Microservices in Payment Systems: A Study of Kubernetes and Docker in Financial Applications. *Universal Journal of Business and Management*, 1(1), 1-17.
- Dong, X. L., Rekatsinas, T., Gokhale, C., & Thirumuruganathan, S. (2020). Data integration and machine learning: A natural synergy. In *Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data* (pp. 2835–2838).
- Bandi, V. D. V. K. Autonomous Data Platforms: Converging AI, MLOps, and Cloud Engineering for Digital.
- Abedjan, Z., Golab, L., & Naumann, F. (2015). Profiling relational data: A survey. *The VLDB Journal*, 24(4), 557–581.
- Chu, X., Ilyas, I. F., Krishnan, S., & Wang, J. (2016). Data cleaning: Overview and emerging challenges. In *Proceedings of the 2016 International Conference on Management of Data* (pp. 2201–2206).
- Amistapuram, K. (2021). Digital Transformation in Insurance: Migrating Enterprise Policy Systems to .NET Core. *Universal Journal of Computer Sciences and Communications*, 1(1), 1-17.
- Rekatsinas, T., Chu, X., Ilyas, I. F., & Ré, C. (2017). HoloClean: Holistic data repairs with probabilistic inference. *Proceedings of the VLDB Endowment*, 10(11), 1190–1201.
- Botlagunta Preethish Nandan. (2021). Enhancing Chip Performance Through Predictive Analytics and Automated Design Verification. *Journal of International Crisis and Risk Communication Research*, 265–285. <https://doi.org/10.63278/jicrcr.vi.3040>.
- Krishnan, S., Franklin, M. J., Goldberg, K., & Wu, E. (2016). ActiveClean: Interactive data cleaning for statistical modeling. *Proceedings of the VLDB Endowment*, 9(12), 948–959.

- Chu, X., Morcos, J., Ilyas, I. F., Ouzzani, M., et al. (2015). KATARA: Reliable data cleaning with knowledge bases and crowdsourcing. *Proceedings of the VLDB Endowment*, 8(12), 1952–1955.
- Kolla, S. H. (2026). Autonomous Enterprise Agents: Orchestrating Large and Small Language Models for Scalable Decision Automation in ITSM, HRSD, and CSM Platforms. *INTERNATIONAL JOURNAL OF ADVANCES IN SIGNAL AND IMAGE SCIENCES*, 24-45.
- Wang, J., Krishnan, S., Franklin, M. J., Goldberg, K., & Ré, C. (2021). Learning to clean data with data programming. *Proceedings of the VLDB Endowment*, 14(11), 2107–2120.
- Ratner, A., Bach, S., Ehrenberg, H., Fries, J., et al. (2017). Snorkel: Rapid training data creation with weak supervision. *Proceedings of the VLDB Endowment*, 11(3), 269–282.
- Uday Surendra Yandamuri. (2023). An Intelligent Analytics Framework Combining Big Data and Machine Learning for Business Forecasting. *International Journal Of Finance*, 36(6), 682-706. <https://doi.org/10.5281/zenodo.18095256>
- Ratner, A., Hancock, B., Dunnmon, J., Sala, F., et al. (2020). Snorkel MeTaL: Weak supervision for multi-task learning. In *Proceedings of The Web Conference 2020* (pp. 2961–2968).
- Jagadish, H. V., Gehrke, J., Labrinidis, A., Papakonstantinou, Y., et al. (2014). Big data and its technical challenges. *Communications of the ACM*, 57(7), 86–94.
- Nagubandi, A. R. (2024). Breakthrough Real-Time AI-Driven Regulatory Intelligence for Multi-Counterparty Derivatives and Collateral Platforms: Autonomous Compliance for IFRS, EMIR, NAIC, SOX & Emerging Regulations. *Journal of Information Systems Engineering and Management*, 9.