

The trade-offs between Monolithic vs. Distributed Architectures

Matheus Felisberto [FIA Business School | matheus.felisberto@gympass.com]

Abstract Software architects frequently engage in trade-off analysis, often confronting sub-optimal solutions due to unforeseen or overlooked disadvantages. Such outcomes can detrimentally affect a company's business operations and resource allocation. This article conducts a critical review of architectural styles, particularly focusing on the strengths and weaknesses of both monolithic and distributed architectures, and their relationship to architectural characteristics. It also explores the role of cloud computing in transitioning from monolithic to distributed-based applications. Utilizing a broad range of sources, including papers and books from both industry and academia, this research provides an overview from theoretical foundations to practical applications. A notable trend observed is a shift back from distributed to monolithic architectures, possibly due to factors such as cost, complexity, and performance.

Keywords: Architectural Styles, Cloud Computing, Microservices, Monoliths, Distributed Transactions, Distributed Architectures

© Published under the Creative Commons Attribution

1 Introduction

As technology evolves in a fast-paced environment, cloud computing has emerged significantly reshaping how software is designed, deployed, and utilized. This shift towards a cloud-native era, favoring agility, scalability, and elasticity, represents a move away from traditional monolithic, on-premises solutions. In conjunction, new architectural styles like serverless computing, have become a viable and popular alternative for many businesses.

However, an often-overlooked aspect is the trade-offs involved. The complexities introduced by these architectural choices are usually challenging to modify, as per the definition of software architecture by [Fowler, 2002]. The trade-offs may range from underutilizing a system, such as a database management systems, to dealing with the intricate complexities of a distributed transaction.

This article aims to provide a systematic review of architectural styles exploring both benefits and inherent challenges between monolithic and distributed-based architectures in what might be described as the no-free-lunch theorem scenario in software design.

2 Methodology

This article employs a methodology that blends snowballing research with a grey literature review, focusing on architectural styles and characteristics between distributed architectures and monoliths. This study began with a search for foundational papers and reports in these fields, using academic databases and industry publications. Through the snowballing process, the base literature collection was progressively expanded. To

ensure a comprehensive coverage, I traced backward and forward the references, capturing both seminal and recent works. Supplementing these academic findings, white papers, technical reports, blog posts and industry case studies were also reviewed to provide a practical industry insights.

The gathered materials were deeply analyzed to identify common themes, trends and divergences in the field. This phase involved extracting and synthesizing key insights related to the research themes, leading to the findings presented in this article. In conjunction snowballing research with grey literature review, enabled a thorough and multifaceted exploration of the topics, capturing both theoretical and practical insights.

3 Software Architecture

Software Architecture plays a crucial role in the industry, however, it still lacks a descriptive and universally convincing definition. Over the years, many experts have attempted to delineate its boundaries, often in the presence of exceptions and caveats. It can be thought of as the composition of a system's components and the relationships between them. The primary focus is on the public interface; the architectural significance lies in these externally facing elements rather than in the private details. However, even though they are not the main concern, can still play a role in shaping the overall structure and functionality of the system, underscores Bass *et al.* [2012].

Building software on a solid foundation, which might include robust architectural patterns and coding standards, is essential, especially considering the key role it plays in a business's success. Failing to address a wide range of potential scenarios could significantly mining a

company's growth. While many modern frameworks offer solutions to common problems, supporting in these architectural decisions, it's important to recognize that each system has its own unique requirements and specifications that must be carefully considered, indicates Microsoft [2009].

3.1 Architectural Characteristics

Every software inherently depends on factors beyond its immediate domain, which architects must consider during the design process. These factors, often referred to as Quality Attributes or non-functional requirements, have been studied at least since 1970s by the software community, portrays Bass *et al.* [2012], and are essential for ensuring the software meets both business and user needs. Extending from low-level code aspects, such as modularity, to operational concerns, including scalability, as detailed in Bass *et al.* [2021] and Richards and Ford [2020]. It is important to recognize that every software characteristics list, is an incomplete, ambiguous, and overlapping list, regardless of the very own ISO [2011]. For instance, Richards and Ford [2020] exemplifies that interoperability and compatibility might be equivalent depending on the system. Bass *et al.* [2012] also point out a particular case of denial-of-service attack. Would it be an aspect of availability, performance, security, or usability? In some extent, we can consider them all. One of the main reasons behind the imprecise definitions and the persisting ambiguity, is the fast-pacing evolution in the industry itself. With many different use cases, companies develop their own terminologies, complicating efforts to establish comprehensive taxonomies for all categories. However, it is recommended to adopt the Domain Driven Design, specifically the ubiquitous-language technique, to minimize misunderstandings. As outlined by Evans [2003], this approach emphasizes consistency in both language and terminology across the whole domain, thereby fostering a better comprehension and clearer communication.

Domain experts should object to terms or structures that are awkward or inadequate to convey domain understanding; developers should watch for ambiguity or inconsistency that will trip up design. Evans [2003]

When it comes to architectural characteristics, it is essential to understand that not every application will require or support every characteristic maximization. It is a decision-making process selecting the aspects that are relevant to solve a problem at hand, while taking into consideration the impact among them. Take, for example, a system where security is crucial. By prioritizing this characteristic, it might cause implications on performance, such as the need for encryption at-rest and at-transit, and secure storage strategies for sensitive data. Attempting to excel in every aspect can lead to the development of a generic architecture, which, in its naively attempt to address a wide range of business problems, may lack focus and efficiency.

SLA (% Uptime)	Unavailability per Year
99%	3 days 15 hours
99.9%	8 hours 45 minutes
99.95%	4 hours 22 minutes
99.99%	52 minutes
99.999%	5 minutes

Table 1. AWS availability tiers

Richards and Ford [2020] advocate for an iterative design process. This approach suggests that implementing smaller, iterative changes is more effective than trying to preemptively address every possible scenario, which can lead to over complexity and unanticipated challenges.

Never shoot for the best architecture, but rather the least worst architecture. Richards and Ford [2020]

3.1.1 Availability

Availability refers to the aspect of software being ready and operational when needed. Given the complexity of systems, and the wide range of potential, and unavoidable issues, maintaining high availability is a critical challenge. This involves not only identifying failures cause by faults but also implementing repair mechanisms, ideally autonomous, that is, without the need for human intervention. Among the difficulties in designing highly available, fault-tolerant systems is accurately predicting potential failure modes and then devising effective strategies to mitigate them. As Bass *et al.* [2012] emphasize, this proactive approach to failure detection and mitigation is key to achieving and maintaining system availability.

To quantify a system's availability, Service Level Agreements (SLAs) are typically established. These agreements specify the guaranteed performance standards and outline the penalties for failing to meet these criteria. To illustrate it, the Table 1 shows the SLAs tiers by the public cloud provider Amazon Web Services [2024].

The challenge lies in not designing with the false and tempting notion of avoiding problems entirely, but rather in considering failure as first-class citizens. By adopting this mindset, it becomes possible to identify the types of failures a system is prone to and to implement appropriate techniques to handle them, as highlighted Bass *et al.* [2012].

3.1.2 Interoperability

The interoperability refers to the ability of a different systems to exchange and effectively interpret information. This concept involves two distinct layers: syntactic interoperability, which is concerned with the structure and format of the data exchange, and semantic interoperability, which relates to the meaning and interpretation of the data. These layers are comprehensively described in Bass *et al.* [2012]. Brownsword *et al.* [2004] further define interoperability as the capacity of

a collection of communicating entities to share specified information, and operate on that information according to an agreed operational semantics. An integral part of achieving interoperability is the process by which systems locate each other and manage their interfaces to facilitate the sharing of information.

3.1.3 Modifiability

Software is a continuous work in progress, whether it is already in production or not. Continuous development, ranging from refactoring and adding new features to updating dependencies, is a norm. This could involve upgrades to the runtime, framework, or even the hardware itself. According to Bass *et al.* [2012], there are two main tactics to improve a system's modifiability: enhancing cohesion and reducing coupling. Improving the former and minimizing the latter can significantly boost this characteristic. A notable contribution to this field is by Dijkstra [1967], who presented the Technische Hogeschool Eindhoven (THE) operating system. This work was an early example of a well-structure operating system, and introduced the hierarchical model, a novel approach at the time. This approach promotes a better modifiability, as each layer is designed with distinct responsibilities, abstraction levels, control flows, and modular structures. In the same vein, Parnas [1972] heavily influenced software design towards modularization, introducing the concept of "information hiding". According to Parnas, modules should have a well-defined interface and hide its working from the rest of the system. This philosophy enables individual modules modifications without affecting others, provided the interfaces remain consistent. Parnas also emphasized that modules are likely to suffer changes throughout their lifecycle, highlighting the importance of reducing inter-module dependencies and hiding information.

3.1.4 Performance

According to the ISO [2011], performance can be understood through three primary aspects: time behavior, resource utilization, and capacity. These elements collectively define how efficiently a software performs a specific computational task within predetermined time constraints and throughput requirements, while also optimizing resource consumption. The field of performance is broad and has been the subject of extensive research. It necessitates a deep understanding of both hardware and software interactions. As Liu [2009] describes, when the performance of one component is pushed to its limits, the other often plays a critical compensating role.

Performance is particularly interesting as an architectural characteristic because it is directly perceived by the user, as much as subsection 3.1.1. Consider, for instance, a trading system suffering by latency or an online game where commands are delayed. Such frustrations can drive users away. Performance can be measured in terms of throughput or latency, depending on the application's requirements. In a web server handling

thousands of requests, achieving high throughput is desirable. Conversely, online gaming demands low latency. Throughput refers to the amount of data a system can process in a given period, whereas latency measures the time it takes for a request to travel from sender to receiver and back. Improving performance might involve reducing demand, such as refusing to serve requests in an overwhelmed web server, or managing resources more effectively through scheduling, replication, or increasing the resource pool, as explained by Bass *et al.* [2012].

3.1.5 Security

Is well known that data has become as valuable as oil in the modern era. In the context of software, security is a critical characteristic that measures a system's ability to protect its data, particularly from unauthorized access. According to Bass *et al.* [2012], there are three main aspects of security: confidentiality, integrity, and availability. Confidentiality ensures that data is accessible only to authorized users. For example, only you should have access to your mobile banking application. Integrity involves safeguarding data from unauthorized modifications ensuring that any alteration is done only through legitimate means. Finally, availability refers to the system being accessible and usable when legitimate users need it.

In the recent years, regulations like the General Data Protection Regulation (GDPR) have been established to address growing privacy concerns. A key aspect of these regulations is determining which information can be shared and with whom. This decision-making process is crucial, not just regarding what information is shared, but also the nature of the information being shared. Consider this scenario: you forget your password and initiate a process to recover your mobile account. You enter your email and set a new password. In the event of a data breach, you can change your email or password, albeit with some inconvenience. But what if the leaked information included your biometric data?

Security is about more than just encrypting data, whether at-rest or in-transit. It also involves meticulously tracking each piece of data, understanding why it's needed, and ensuring traceability, as demanded by regulatory bodies. This comprehensive approach to security helps protect sensitive information, especially in cases where it cannot be easily changed or replaced, like biometric data.

3.1.6 Testability

To ensure that software is functioning correctly and ideally free of known faults or bugs, it's important to understand what constitutes a software bug. Patton [2005] outlines at least five scenarios that can be considered as software bugs:

1. *The software fails to do something that the product specification states it should do.*

2. *The software does something that the product specification explicitly says it should not do.*
3. *The software performs an action that the product specification does not cover.*
4. *The software omits an action that is not mentioned in the product specification but is implied or necessary.*
5. *The software is user-unfriendly, slow, or appears flawed to the software tester or the end user.*

Testability, the ease with which software can be tested, is crucial in this context. It refers to how readily a software system, either in isolation or in combination, can be tested for bugs. This is often facilitated through test harness tools that control inputs, outputs, and possibly the state of the software, ensuring it aligns with the provided specifications. High testability is linked to increased reliability, making it a valuable characteristic to consider in software architecture.

3.1.7 Usability

Usability is a measure of how easily users can interact with a system. Its significance has increased over the years, becoming a critical factor in users' decision-making when choosing between services. Bass *et al.* [2012] describes several dimensions of usability. One key aspect is the ease of learning system features, often referred to as onboarding. For example, how quickly can a new user learn to use the system? Implementing short tutorials is a useful tactic in this regard. Efficient use of the system and minimizing the impact of errors are also crucial. The former could involve allowing users to pause and later resume a long-running task. The latter ranges from displaying errors in a user-friendly manner to providing recovery options, like retrying or suggesting alternative steps to achieve the same objective. Adapting the system to meet user needs is another way to enhance usability. A familiar example is a browser auto-filling credit card details based on previous usage. Lastly, boosting users confidence and satisfaction also plays a role in enhancing usability. For instance, providing visual feedback during a video upload, like showing a progress percentage, can significantly improve the user experience.

3.2 Architectural Styles

Architectural styles represent the macro structure of a system and are categorized based on a set of characteristics, each offering distinct advantages and disadvantages. One of the key responsibilities of an architect is to select the appropriate architectural style by evaluating it against the business requirements. The following sections systematically describe both monolithic and distributed-based architectural styles in a non-exhaustive manner.

3.2.1 Layered Architecture

The layered architecture is one of the most straightforward and commonly used architectural styles, appreciated for its simplicity, popularity, and cost-effective implementation, as described by Richards and Ford [2020]. This architectural style naturally aligns with the needs and structures of many traditional business applications.

The topology of a layered architecture style, as illustrated in Figure 1, is characterized by logical horizontal layers, each assuming a specific role within the application. While there is no strict limitation on the number of layers, four standard ones are commonly observed: presentation, business, persistence, and database. The actual number and nature of these layers can vary depending on the size and complexity of the application.

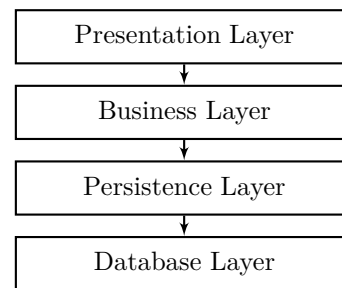


Figure 1. Layered-architecture style topology

Each layer in a layered architecture has a distinct role and responsibility. For example, the business layer handles the application's business logic, while the presentation layer manages browser communication and interface logic. This structure ensures that each layer abstracts its specific function to collectively fulfill a business request. For instance, the presentation layer doesn't concern itself with customer details, which are the purview of the business layer, nor does the business layer need to manage HTML markup. This separation of concerns facilitates the construction of layered architectures with clearly defined responsibilities for each layer. However, as Richards and Ford [2020] point out, one drawback of this approach is the challenge in applying domain-driven design. In layered architectures, the business domain tends to be dispersed across all layers, categorized by technical role rather than being grouped into components like Customer.

Suggested applications The layered architecture is a great choice in the early stages of a project, particularly when there is significant uncertainty about the product itself but development needs to begin. As Richards [2022] details, it is also well-suited for projects with time or budget constraints. Being a monolithic style, it avoids the complexities associated with distributed systems, such as the need to define contracts and manage remote procedure calls. These aspects of distributed systems will be discussed in more detail in Section 4.1.

Discouraged applications One important consideration with layered architecture is its scalability challenge. As a monolithic architectural style, scaling up for improved performance, elasticity, or overall scalability typically involves scaling the entire application. This approach may not always be efficient. Furthermore, layered architectures may lack sufficient fault tolerance. Since all components are deployed together, a fault in one area can potentially bring down the entire system, as Richards [2022] notes.

3.2.2 Microkernel architecture

Inspired by operating systems, the microkernel architecture style is known for its flexibility and extensibility through the use of plugins, often referred to as plug-in architecture. Although it was developed several years ago, it remains widely used today. Its extensibility lies in maintaining core functionality while allowing for the addition of varied features through plugins. This architecture comprises two fundamental components: the core and plugin modules. The core provides essential functionalities, whereas plugin modules add additional, usually standalone, independent components. A common method for the core system to recognize available plugins is through a plugin registry, which contains all necessary information about them. While these plugins have traditionally been developed as libraries, they can also be implemented through remote services such as REST, according to Richards [2022].

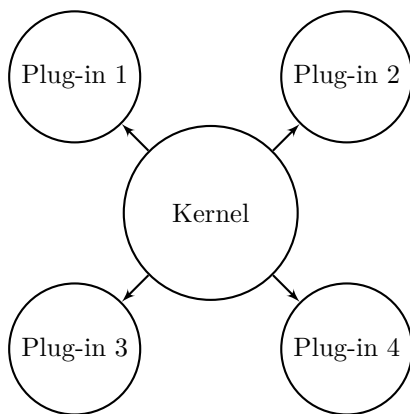


Figure 2. Microkernel architecture style topology

Suggested applications This architectural style is an excellent choice for products that require customization or will be expanded with more features over time, providing an easy starting point for such projects. The microkernel architecture also adapts well to varying configurations or deployment requirements. For instance, a microkernel application can be deployed in an on-premise facility with a custom set of plugins tailored to specific needs, while maintaining core functionality that remains unchanged and completely agnostic. Richards [2022] emphasizes that the microkernel, as well as discussed in subsection 3.2.1, offers a cost-effective and

relatively simple setup, making it a practical option for projects with limited time or budget.

Discouraged applications The downsides of this architectural style lie in its heavy reliance on the core module, which can act as a bottleneck. This dependency complicates elasticity, reduces fault tolerance, and is less suited for highly scalable solutions. Additionally, the architecture discourages frequent modifications to the core module, preferring instead to use plugins to add features and extend functionality. When frequent core updates are necessary, this approach becomes inadequate. As Richards [2022] suggests, if adhering to this rationale is not feasible, this architectural style may not be the best solution to the problem at hand.

3.2.3 Event-Driven architecture

Event-driven architecture is a distributed, asynchronous architectural style that has become increasingly popular, especially following the democratization of the cloud by public vendors. This architecture excels in managing complex workflows and fostering reactive, responsive systems. It has also boosted the creation of tools, frameworks, and cloud-based solutions tailored to enhance accessibility and ease of setup.

In an event-driven system, decoupled event processing components receive and process events asynchronously. The process is initiated by an event, which kicks off an asynchronous workflow. For example, consider placing an order on an e-commerce platform that uses event-driven architecture. The moment you place the order, an event is created containing all necessary details. This event triggers the asynchronous workflow. Subsequently, a component, such as the payment processor, receives this event, processes it, and may generate another event to inform the system of a state change. Once the payment is confirmed, another event is published, prompting the shipment component to act. All these interactions involve physical messaging artifacts. Typically, initiating events utilize point-to-point channels through messaging services or queues, while processing events often employ publish-subscribe channels via topics or notification services, as described by Richards [2022]. Event-driven architecture can function as a standalone architectural style or be combined with others, such as microservices.

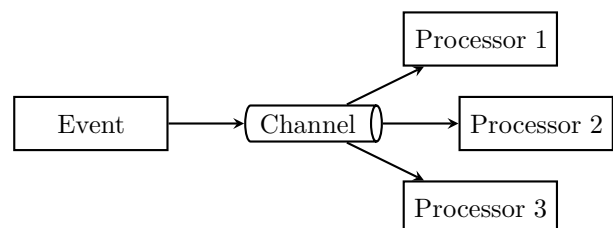


Figure 3. Event-driven architecture style topology

Suggested applications Event-driven architecture is undoubtedly an optimal choice for achieving high performance, scalability, and fault tolerance. It also brings additional benefits as it shifts the paradigm from traditional decision-tree logic to a model that reacts to events. This approach facilitates the modeling of highly complex systems by isolating each component within its context, focusing on its specific responsibilities, and processing incoming events. Moreover, it can aid organizational scaling by allowing teams to work independently on their components, each defined by clear contracts. However, the discussion of team typologies is out of scope of this article.

Discouraged applications By its nature, being a distributed asynchronous architecture, the event-driven architectural style is not recommended for request-based systems that require high data consistency, as revealed by Richards [2022]. Due to its asynchronous nature and the eventual consistency of processing, meaning there is no guarantee of when an event will be processed, event-driven architecture may not suit scenarios demanding immediate data accuracy, such as CRUD operations.

Additionally, managing the order in which events are consumed and handling errors present significant challenges. Sequencing issues can complicate the architecture, especially if events need to be combined to produce further actions, or the potentially deadlock situations where, for instance, event A waits for event B, event B waits for event C, and event C waits for event A. Error handling adds another layer of complexity. For example, consider a scenario where a payment event is processed successfully, but a subsequent inventory event fails because the item is out of stock. In such cases, stakeholders must clearly define responsibilities and corrective actions for each potential issue, which can lead to corner cases that compromise the system's reliability in terms of data consistency.

3.2.4 Microservice architecture

Microservice architecture is arguably the most popular and widely used architectural style today. The term 'microservice' was coined in 2014 by Martin Fowler and James Lewis. Although they were the first to publish an article detailing its characteristics, they emphasized that this architectural style was not a novelty but rather rooted in UNIX philosophy. Additionally, microservices incorporate Domain-Driven Design concepts, such as bounded contexts, which organize entities and behaviors within well-defined limits in code and database schemas. Microservices operate by integrating multiple smaller, independently deployed services that communicate with each other. This approach is defined by Fowler and Lewis as follows:

In short, the microservice architectural style is an approach to developing a single application as a suite of small services, each running in its own process and communicating

with lightweight mechanisms, often an HTTP resource API. These services are built around business capabilities and independently deployable by fully automated deployment machinery. Fowler and Lewis [2014]

The primary goal of microservice architecture is to achieve decoupling, with physical modeling that represents bounded contexts. To accomplish this, each service is designed to be self-contained, possessing all necessary components for operation, including its own database systems. This approach eliminates a single source of truth, allowing different services to utilize diverse technologies tailored to their specific needs. For example, in contrast to a monolithic architecture like a subsection 3.2.1, where performing a full-text search using a traditional RDBMS might be unsatisfactory at scale, a microservice architecture could leverage a specialized Elasticsearch or Solr cluster to handle such tasks more effectively.

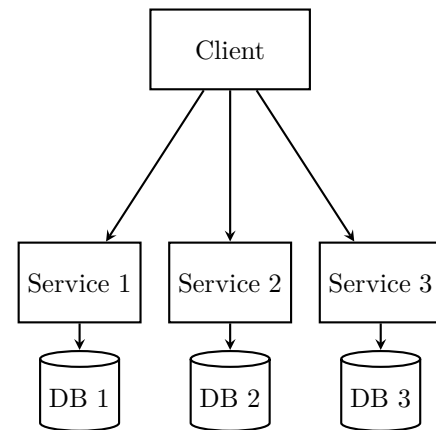


Figure 4. Microservice architecture style topology

As a distributed architecture, each service in a microservice setup must run independently in its own process. A significant factor in the widespread adoption of this architectural style has been the rise of cloud computing, which has made managing these physical constraints easier through virtualization and containers. According to Richards and Ford [2020], decoupling services to the extent that each domain maintains its own infrastructure was once impractical. However, with the availability of free open-source operating systems and automated infrastructure management, it is now feasible at both the domain and operational levels.

Network calls present a bottleneck in this architecture, as they are typically slower than straightforward method calls due to the process of transferring data across a network, which often includes a layer of security. Additionally, communication between microservices requires that each service knows how to initiate contact with others without relying on a centralized orchestrator, to prevent coupling. Richards and Ford [2020] characterizes this communication as protocol-aware, meaning the caller must know the communication protocol in advance, and heterogeneous, allowing services to be

developed using different technologies, thus ensuring interoperability.

Suggested applications When elasticity, high fault tolerance, and scalability are priorities, the microservice architecture style is an excellent choice. Extensibility is another significant advantage; adding functionality typically involves spinning up new services within the already standardized infrastructure. This flexibility is crucial, allowing rapid adaptations and changes aligned with business needs. The architecture also enhances maintainability due to each service's smaller scope and the clear boundaries guaranteed by the concept of bounded contexts. Testing is generally more straightforward, as long as dependencies on other services are minimized. When dependencies are unavoidable, a robust mocking strategy can effectively facilitate testing.

Discouraged applications The primary goal of microservices is decoupling. However, this architecture may not be suitable in scenarios where, despite physical decoupling, there remains substantial logical coupling. This is particularly problematic in systems with complex workflows and excessive inter-service communication. Often, the need for frequent communication between services is driven by data dependencies, which are a critical factor in deciding whether to adopt this architectural style. The more tightly coupled the data, the less effective the microservice architecture becomes.

Microservices represent one of the most complex architectures in use today. They can also be costly due to the duplication of efforts, including data replication across multiple databases. As the number of services increases, the infrastructure required for each service also grows exponentially, as delineated by Richards and Ford [2020].

Despite what may seem counter intuitive, microservices are not inherently high-performance or highly-responsive systems. Their performance is often hampered by three types of latency: network latency, security checks, and data access delays. Each of these factors needs to be fine tuned to the specific needs of each service, even at the database level.

4 Architectural decisions

One of the key responsibilities of a software architect is to analyze a range of factors that differ from business to business and technology to technology. Unfortunately, there is no one-size-fits-all solution. Choosing an architectural style requires careful consideration of trade-offs, ensuring that the disadvantages are not overlooked in favor of the advantages. Many factors contribute to this necessity. One significant factor is the hype surrounding new frameworks, language features, or even new programming languages. It is natural for developers and technologists to be excited by the purported

innovations these new tools bring. However, this enthusiasm can sometimes lead to solutions that are overly complicated, costly, and complex. Ford *et al.* [2021] describe this phenomenon as evangelism, where the focus on the positive aspects leads to neglecting potential downsides.

In order to mitigate risks while ensuring that the architecture evolves in line with business requirements, it is crucial to analyze and prioritize what is most important at any given time, or to project future needs as discussed in section 3.1. Unfortunately, distributed architectures, particularly microservices, are often default choices for many applications. This trend is driven by misconceptions and an underestimation of the trade-offs involved. The widespread popularity of this architectural style is not coincidental, and it necessitates a careful consideration of its implications.

Recently, even big tech companies, equipped with their own data centers and substantial intellectual capital, have begun shifting toward monolithic architectures. This move aims to reduce costs, complexity, and scalability issues. Ghemawat *et al.* [2023] propose a programming methodology that reportedly reduces latency by 15 times and costs by 9 times compared to the current standard. The paper also critiques some commonly touted benefits of microservices, pointing out how they can actually lead to problems in performance, correctness, and development agility. For example, performance can suffer when the granularity of services is poor, leading to excessive inter-service communication. Correctness becomes challenging when considering all possible interactions between services, including variations in service versions and handling failures. Development agility is compromised when changes across services cannot be made atomically, requiring extensive coordination effort. To address these issues, the proposed methodology advocates for writing monolithic applications that are modularized as components, assigning these components to physical processes, and deploying them atomically.

Another compelling example of cost reduction achieved by shifting from a distributed microservices-based architecture to a monolith is the Amazon Prime Video case. This case study describes how the system was initially orchestrated by a component that not only created a bottleneck but also incurred the highest costs within the solution. After transitioning to a monolithic architecture, the infrastructure costs were reduced by an astonishing 90%, and it is also reported to have enhanced the system's scalability, as described by Kolny [2023].

4.1 Fallacies of distributed systems

A common misconception about transitioning from a monolithic to a distributed architecture is encapsulated in the 8 fallacies of distributed systems. For instance, a local method call used to process a business rule is not equivalent to inter-service communication that relies on the network. Numerous precautions are neces-

sary in such environments, including robust error handling. While TCP is a reliable protocol that will resend any lost packets, higher levels in the stack, such as those involving REST APIs, still require individual request management. Additionally, latency increases are a concern. Even though the TCP slow start can be mitigated by establishing long-lived connections, the overhead from serialization and deserialization, along with additional security checks, can significantly extend the final processing time.

1. **The network is reliable:** Assumes that network connections are always stable and reliable.
2. **Latency is zero:** Ignores the time delay in communication over a network.
3. **Bandwidth is infinite:** Ignores the limitations in network data transfer capacity.
4. **The network is secure:** Assumes that the network is naturally secure from attacks and vulnerabilities.
5. **Topology does not change:** Assumes that the way network systems are connected remains constant.
6. **There is one administrator:** Ignores the complexity of managing distributed systems with multiple administrative domains.
7. **Transport cost is zero:** Ignores the resources and time required to move data across the network.
8. **The network is homogeneous:** Assumes that the network's hardware, software, and protocols are consistent and compatible.

4.2 Data management

Data management is undoubtedly one of the most complex and critical aspects to consider when comparing monolithic and distributed architectures. In a monolithic architecture, data management is generally simpler because the same system both writes to and reads from a single database. This setup naturally supports transactional capabilities, ensuring that ACID (Atomicity, Consistency, Isolation, Durability) properties are maintained. There is less concern about data integrity, consistency, and security since authentication and authorization are more straightforward in a unified environment. Querying data is also relatively simple, although it may require complex join clauses that necessitate the creation of views. Nevertheless, this setup is typically sufficient for the needs of most small to mid-sized businesses.

Conversely, the more distributed an architecture becomes, the more challenging it is to maintain these properties. In such environments, ensuring data integrity, consistency, and security across multiple services and databases can be daunting. It is usually inadvisable — except in exceptional cases, to attempt to manage these aspects manually due to the increased complexity and risk involved.

4.2.1 Transactional complexity

Transactions are a fundamental aspect of enterprise applications, with ACID properties helping to shield developers from many concerns related to data integrity. In monolithic architectures, these properties are typically leveraged without additional complications, as the architecture does not introduce inherent challenges in managing data consistency. However, distributed architectures, particularly those utilizing a database-per-service pattern, can encounter difficulties when transactions need to update data distributed across multiple services.

To address this, the Saga Pattern is employed, which orchestrates a sequence of local, message-driven transactions to ensure data consistency across services, as detailed in Richardson [2018]. Yet, implementing the Saga Pattern introduces its own challenges, such as no isolation which necessitates countermeasures to manage concurrency issues effectively. Sagas can be managed through two approaches: choreography, where each local transaction publishes events that trigger subsequent transactions in other services, and orchestration, where a central coordinator directs each participant to execute their transaction. In cases of failure, the saga must initiate compensating transactions in reverse order to undo the effects, as illustrated in Richardson [2018].

4.2.2 Eventual Consistency

To enhance architectural characteristics such as availability and performance, adopting eventual consistency is a strategic decision. Ford *et al.* [2021] identifies three main patterns that facilitate this approach: background synchronization, orchestrated request-based pattern, and event-based pattern.

Background Synchronization Pattern This pattern employs an external system to periodically check and update all data sources to ensure they remain synchronized. The level of eventual consistency achieved depends on the frequency of the updates, which could be scheduled as nightly batch jobs or on an hourly basis. Although this approach effectively keeps data in sync, it also tends to couple related services together. This coupling occurs because the pattern requires knowledge of all data schemas, which can lead to the violation of bounded contexts, duplication of business logic, and a complex implementation.

Orchestrated Request-Based Pattern Unlike the Background Synchronization Pattern, the goal of this pattern is to process the entire distributed transaction within the scope of the business request itself. It can be implemented either directly within a specific service or through a dedicated orchestrator service. Implementing it within a service tends to overload that service's responsibilities, as it must manage business rules and handle errors in addition to its regular duties. The key

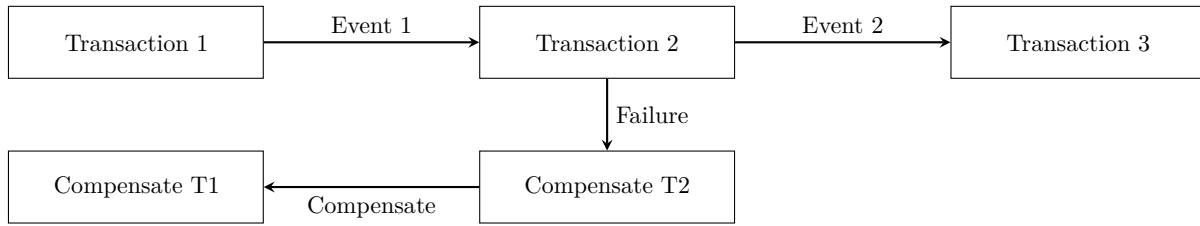


Figure 5. Saga pattern topology

issue with this pattern as an eventual consistency approach is its error handling, if an error occurs during the transaction or during a compensating transaction, there will be no external system available to recover the state.

Event-Based Pattern This pattern utilizes asynchronous publish/subscribe or messaging systems to post events or commands. In this approach, all services are decoupled, and the eventual consistency time factor is minimized due to the asynchronous and parallel nature of the messaging systems, as indicated in Ford *et al.* [2021]. The main challenges, similar to other patterns, revolve around error handling. While the messaging system does not require the consumer to be available at the time the message is published — ensuring they receive it when they are available, failures during the processing of these events can still lead to consistency issues.

4.2.3 Querying data

Monolithic architectures typically rely on a single database, which generally simplifies the task of querying data. However, in distributed applications, querying becomes significantly more challenging due to the complexities involved in managing data across multiple services. Ford *et al.* [2021] outlines four patterns specifically designed to address these challenges.

Interservice Communication Pattern This pattern is commonly utilized to access data from another service, typically through REST APIs or remote procedure calls such as gRPC. However, it introduces several downsides, including those mentioned in subsection 4.1, and negatively impacts the availability due to service coupling. Additionally, this pattern often requires complex error handling strategies. For example, if a request from service A to service B fails due to a transient error, it directly affects the user experience at that moment.

To mitigate such issues, the Retry Pattern, as defined in Azure [2023b], can be implemented. This pattern automatically retries recoverable failures a specified number of times with exponential backoff, enhancing usability and availability. However, it also risks exacerbating issues in the event of non-recoverable errors, such as a database outage, potentially creating a cascading effect within the service chain.

To address this, the Circuit Breaker Pattern, explained in Azure [2023a], prevents the application from

performing operations likely to fail. This safety mechanism can be crucial in maintaining system stability and preventing further complications. Both the Retry and Circuit Breaker patterns can be implemented directly within the application or through the infrastructure using service meshes like Istio and Traefik.

Column Schema Replication Pattern This pattern involves replicating columns to make data available across different bounded contexts, as discussed in Ford *et al.* [2021]. On one hand, this strategy enhances performance and eliminates dependencies on other services when accessing data. On the other hand, maintaining data consistency and ensuring synchronization across services can pose significant challenges. Furthermore, implementing the replication mechanism — whether synchronous or asynchronous, requires careful consideration. The asynchronous approach is often preferable as it enhances system responsiveness and reduces availability issues by minimizing dependencies among services.

Replicated Caching Pattern While caching generally improves performance by allowing data access from memory rather than disk, it can also facilitate access to distributed data. In this pattern, if service A needs to access data from service B, it retrieves this data directly from a cache that contains all data from service B, which are written to the cache upon changes and then replicated. This approach significantly enhances scalability and fault tolerance. However, it involves trade-offs, including service dependency. For instance, as the data owner, service B must populate the cache before service A can begin operations. Additionally, managing large volumes of data can be problematic; according to Ford *et al.* [2021], 500 MB is considered a threshold. Architects must carefully consider the volume of data, cache size, and the number of service instances that require the cache. Furthermore, cache replication might experience delays in synchronizing data between services, making this pattern less suitable for data that changes frequently but more appropriate for data that changes less often. Finally, configuring and managing this pattern are complex tasks that require the services to be aware of each other.

Data Domain Pattern Although sharing data is generally not recommended, this pattern offers benefits that may be worthwhile depending on the context, such as decoupled services which enhance availability

as an architectural characteristic. Responsiveness is another significant advantage, as it eliminates the complications associated with the Interservice Communication Pattern; for example, the required data is just an SQL query away. This pattern also maintains data integrity and consistency, and allows for the use of RDBMS features such as views, procedures, and triggers.

Contrary to the Interservice Communication Pattern and Replicated Caching Pattern, the Data Domain Pattern does not involve contracts or abstract layers over the database, which makes changing database schemas more challenging as it requires coordination across all services within the bounded context. Additionally, there are heightened security and auditing risks, given that services have access to the entire database.

4.3 Security and Auditing

Security is a fundamental and non-negotiable architectural characteristic, especially given the millions of transactions that occur daily on the internet, many of which are fraudulent or non-secure. Both architectural styles must implement security measures, though the complexity and effort involved can vary. In monolithic architectures, there are relatively few entry points into the application, and once within the application layer, there are no additional security measures required among components to enable communication. However, Dias and Siriwardena [2020] highlights that the attack surface is broader in microservice-based architectures, thereby increasing the risk. Distributed security can also impact performance due to the need to connect with a remote security token service, for instance.

Sharing user context is another concern; in a monolithic style, all components share the same session, whereas distributed architectures often use JWTs — a popular choice, to share user context among services, as described by Dias and Siriwardena [2020].

Furthermore, data security and auditing are critical. Encryption is required both in transit and at rest. Distributed architectures naturally have more components to authenticate, authorize, encrypt, and so on, increasing the security complexity.

4.4 Coordinating distributed operations

In contrast to monolithic, distributed-based styles requires a mechanism that enables each component to execute its task. This mechanism can be either an orchestrator or a choreography among the services themselves.

Orchestration Style An orchestrator is a component responsible for managing the workflow state, error handling, and notifications within a domain. Ford *et al.* [2021] makes an important point: microservices-based architectures typically use an orchestrator for each workflow because relying on a single orchestrator, such as in an Enterprise Service Bus, is discouraged due

to the potential for undesirable coupling. The Orchestration Pattern offers excellent error handling, centralized workflow management, recovery capabilities, and state management. However, like any other pattern, it has its disadvantages. Responsiveness can be compromised since all requests must pass through the orchestrator, potentially creating a bottleneck—as was observed in the Amazon Prime case detailed in Kolny [2023]. Moreover, fault tolerance is a concern because the orchestrator itself can become a single point of failure. Scalability can also be challenging due to the orchestrator's numerous coordination points, and the pattern inherently increases service coupling.

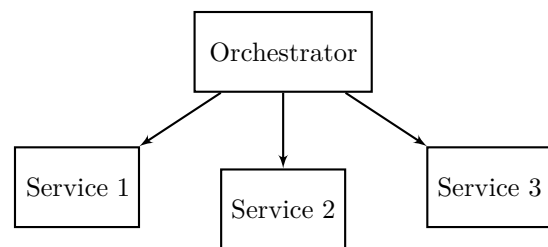


Figure 6. Orchestration style topology

Choreography Style In a choreography style, there is no central orchestrating component; instead, services communicate directly with each other. This approach enhances responsiveness and scalability by reducing the number of intermediaries in the workflow and potentially increasing parallelism. Fault tolerance is also improved due to the ability to scale services independently of orchestrator, which also reduces coupling by the same reason. However, distributing the workflow introduces challenges, particularly in error handling. Services must possess the knowledge typically centralized in an orchestrator, complicating both error resolution and efforts to enhance recoverability. Additionally, this style lacks centralized state management for ongoing operations, as delineated in Ford *et al.* [2021].

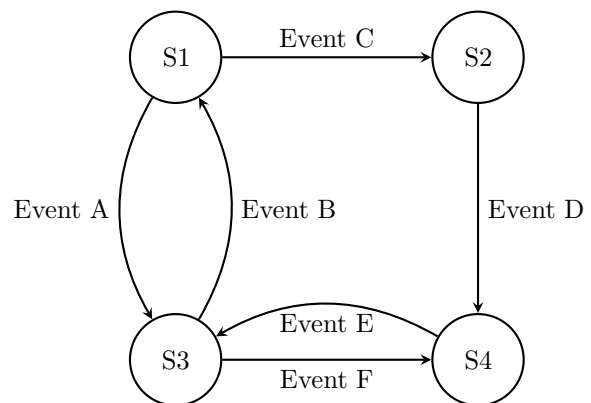


Figure 7. Choreography style topology

4.5 Cost

The widespread acceptance and popularization of distributed architectures were significantly propelled by

the advent of cloud computing. The ease of deploying a virtual machine or container with just a few clicks — or even automatically through scriptable configurations, made it irresistible to explore the vast array of cloud products available. This includes everything from message brokers and various types of databases to load balancers, all accessible to businesses of any size. Thanks to the cloud's elasticity model, billing is based only on the resources actually used. However, costs are incurred in every aspect of usage, including data transfer, storage, access, requests, communication, and authorization, calculated down to the bit level.

The Amazon Prime case revealed by Kolny [2023], documents a dramatic 90% cost reduction achieved by shifting back to a monolithic architecture. According to Su *et al.* [2023], cost is the most common reason companies revert from microservices to monoliths. Operational expenses become unsustainable at a certain scale as each service instance may require its own server (or container), maintain a database, and generate separate metrics and logs.

4.6 Observability

Regardless of the architectural style employed, observability is critical for any serious enterprise application. Merely collecting logs, metrics, and traces is not sufficient, though these are foundational for testing and debugging systems. In Sridharan [2018], observability is described as a property of a system that acknowledges the following realities:

1. *No complex system is ever fully healthy.*
2. *Distributed systems are pathologically unpredictable.*
3. *It's impossible to predict the myriad states of partial failure various parts of the system might end up in.*
4. *Failure needs to be embraced at every phase, from system design to implementation, testing, deployment, and, finally, operation.*
5. *Ease of debugging is a cornerstone for the maintenance and evolution of robust systems.*

In Beyer *et al.* [2016] the Four Golden Signals of monitoring are defined as latency, traffic, errors, and saturation. The complexity of monitoring increases linearly with the number of services and exponentially with the addition of components such as data synchronization tools. Moreover, if you implement patterns like Retry and Circuit Breaker within a Sidecar Pattern - which isolates components in containers, the complexity continues to rise.

As systems become more complex, observability becomes increasingly challenging. In distributed systems, tracing represents a major hurdle. It involves tracking a sequence of causally related events that describe the flow of a request through the system, represented as a directed acyclic graph of spans, as Sridharan [2018] delineates. While one benefit of a distributed architecture is the ability to support polyglot systems, which re-

lies in the interoperability among diverse systems, this same diversity makes tracing more difficult. The challenge arises from the need to instrument a variety of frameworks, languages, and libraries.

5 Conclusion

There is always a trade-off involved with every decision. When an architect encounters a situation where no trade-off seems apparent, it likely has not yet been revealed. An architect's job entails remaining unbiased by the advantages of any particular architectural style or pattern and not aiming solely to enhance all characteristics. Instead, architects must be keenly aware of which disadvantages could be prohibitive due to business constraints. A clear evaluation process is essential, one that distinguishes what is critical, what is fundamental, and what is merely nice to have.

Recently, there has been a trend towards reverting to monolithic architectures due to concerns over costs, complexity, and performance, among other reasons. However, just as the previous enthusiasm for shifting towards microservices should have met with caution, the current trend of moving back to monoliths should also be approached carefully. There is no one-size-fits-all solution, as businesses vary and have their own particularities. Choosing an architectural style is not a question of right or wrong; styles can even be combined to achieve specific goals. It's important to understand the business needs and consider the foreseeable future, while remaining grounded in reality.

Research methodologies that support decisions by evaluating business constraints in contrast to trade-offs among architectural styles are crucial. Additionally, investigating how to improve these architectural styles by combining the strengths of one with the weaknesses of another is fundamental to the evolution of software engineering.

References

- Amazon Web Services, A. (2024). Availability. [Online; accessed 23-March-2024].
- Azure (2023a). Circuit breaker pattern. [Online; accessed 18-April-2024].
- Azure (2023b). Retry pattern. [Online; accessed 18-April-2024].
- Bass, L., Clements, P., and Kazman, R. (2012). *Software Architecture in Practice, Third Edition*. Addison-Wesley Professional.
- Bass, L., Clements, P., and Kazman, R. (2021). *Software Architecture in Practice, 4th Edition*. Addison-Wesley Professional.
- Beyer, B., Jones, C., Murphy, N. R., and Petoff, J. (2016). *Site Reliability Engineering*. O'Reilly Media, Inc.
- Brownsword, L., Carney, D. J., Fisher, D., Lewis, G., Morris, E. J., Place, P. R., Smith, J., Wrage, L., ,

- and Meyers, B. C. (2004). Current perspectives on interoperability. *Software Engineering Institute*.
- Dias, W. K. A. N. and Siriwardena, P. (2020). *Microservices Security in Action*. Manning Publications.
- Dijkstra, E. W. (1967). The structure of the “the”-multiprogramming system.
- Evans, E. (2003). *Domain-Driven Design: Tackling Complexity in the Heart of Software*. Addison-Wesley Professional.
- Ford, N., Richards, M., Sadalage, P., and Dehghani, Z. (2021). *Software Architecture: The Hard Parts*. O’Reilly Media, Inc.
- Fowler, M. (2002). *Patterns of Enterprise Application Architecture*. Addison-Wesley Professional.
- Fowler, M. and Lewis, J. (2014). Microservices.
- Ghemawat, S., Grandl, R., Petrovic, S., Patel, M. W. P., Posva, I., and Vahdat, A. (2023). Towards modern development of cloud applications.
- ISO (2011). Iso/iec 25010. [Online; accessed 22-March-2024].
- Kolny, M. (2023). Scaling up the prime video audio/video monitoring service and reducing costs by 90 [Online; accessed 15-April-2024].
- Liu, H. H. (2009). *Software Performance and Scalability: A Quantitative Approach*. Wiley-Blackwell.
- Microsoft (2009). *Microsoft® Application Architecture Guide, 2nd Edition (Patterns & Practices)*. Microsoft Press.
- Parnas, D. L. (1972). On the criteria to be used in decomposing systems into modules.
- Patton, R. (2005). *Software Testing, Second Edition*. Sams.
- Richards, M. (2022). Software architecture patterns, 2nd edition. Technical report, O’Reilly Media, Inc.
- Richards, M. and Ford, N. (2020). *Fundamentals of Software Architecture*. O’Reilly.
- Richardson, C. (2018). *Microservices Patterns*. O’Reilly Media, Inc.
- Sridharan, C. (2018). *Distributed Systems Observability*. O’Reilly Media, Inc.
- Su, R., Li, X., and Taibi, D. (2023). Back to the future: From microservice to monolith.