

CLOUD ARCHITECTURES, NETWORKS, SERVICES, AND MANAGEMENT

Raouf Boutaba¹ and Nelson L. S. da Fonseca²

¹*D.R. Cheriton School of Computer Science, University of Waterloo, Waterloo,
Ontario, Canada*

²*Institute of Computing, State University of Campinas, Campinas,
São Paulo, Brazil*

1.1 INTRODUCTION

With the wide availability of high-bandwidth, low-latency network connectivity, the Internet has enabled the delivery of rich services such as social networking, content delivery, and e-commerce at unprecedented scales. This technological trend has led to the development of cloud computing, a paradigm that harnesses the massive capacities of data centers to support the delivery of online services in a cost-effective manner. In a cloud computing environment, the traditional role of service providers is divided into two: *cloud providers* who own the physical data center and lease resources (e.g., virtual machines or VMs) to service providers; and *service providers* who use resources leased by cloud providers to execute applications. By leveraging the economies-of-scale of data centers, cloud computing can provide significant reduction in operational expenditure. At the same time, it also supports new applications such as big-data analytics (e.g., MapReduce [1]) that process massive volumes of data in a scalable and efficient fashion. The rise of cloud computing has made a profound impact on the development of the IT industry in recent years. While large companies like Google, Amazon, Facebook,

Cloud Services, Networking, and Management, First Edition.

Edited by Nelson L. S. da Fonseca and Raouf Boutaba.

© 2015 John Wiley & Sons, Inc. Published 2015 by John Wiley & Sons, Inc.

and Microsoft have developed their own cloud platforms and technologies, many small companies are also embracing cloud computing by leveraging open-source software and deploying services in public clouds.

However, despite the wide adoption of cloud computing in the industry, the current cloud technologies are still far from unleashing their full potential. In fact, cloud computing was known as a buzzword for several years, and many IT companies were uncertain about how to make successful investment in cloud computing. Fortunately, with the significant attraction from both industry and academia, cloud computing is evolving rapidly, with advancements in almost all aspects, ranging from data center architectural design, scheduling and resource management, server and network virtualization, data storage, programming frameworks, energy management, pricing, service connectivity to security, and privacy.

The goal of this chapter is to provide a general introduction to cloud networking, services, and management. We first provide an overview of cloud computing, describing its key driving forces, characteristics and enabling technologies. Then, we focus on the different characteristics of cloud computing systems and key research challenges that are covered in the subsequent 14 chapters of this book. Specifically, the chapters delve into several topics related to cloud services, networking and management including virtualization and software-defined network technologies, intra- and inter- data center network architectures, resource, performance and energy management in the cloud, survivability, fault tolerance and security, mobile cloud computing, and cloud applications notably big data, scientific, and multimedia applications.

1.2 PART I: INTRODUCTION TO CLOUD COMPUTING

1.2.1 What Is Cloud Computing?

Despite being widely used in different contexts, a precise definition of cloud computing is rather elusive. In the past, there were dozens of attempts trying to provide an accurate yet concise definition of cloud computing [2]. However, most of the proposed definitions only focus on particular aspects of cloud computing, such as the business model and technology (e.g., virtualization) used in cloud environments. Due to lack of consensus on how to define cloud computing, for years cloud computing was considered a buzz word or a marketing hype in order to get businesses to invest more in their IT infrastructures. The National Institute of Standards and Technology (NIST) provided a relatively standard and widely accepted definition of cloud computing as follows: “cloud computing is a model for enabling ubiquitous, convenient, on-demand network access to a shared pool of configurable computing resources (e.g., networks, servers, storage, applications, and services) that can be rapidly provisioned and released with minimal management effort or service provider interaction.” [3]

NIST further defined five essential characteristics, three service models, and four deployment models, for cloud computing. The five essential characteristics include the following:

1. On-demand self-service, which states that a consumer (e.g., a service provider) can acquire resources based on service demand;

2. Broad network access, which states that cloud services can be accessed remotely from heterogeneous client platforms (e.g., mobile phones);
3. Resource pooling, where resources are pooled and shared by consumers in a multi-tenant fashion;
4. Rapid elasticity, which states that cloud resources can be rapidly provisioned and released with minimal human involvement;
5. Measured service, which states that resources are controlled (and possibly priced) by leveraging a metering capability (e.g., pay-per-use) that is appropriate to the type of the service.

These characteristics provide a relatively accurate picture of what cloud computing systems should look like. It should be mentioned that not every cloud computing system exhibits all five characteristics listed earlier. For example, in a private cloud, where the service provider owns the physical data center, the metering capability may not be necessary because there is no need to limit resource usage of the service unless it is reaching data center capacity limits. However, despite the definition and aforementioned characteristics, cloud computing can still be realized in a large number of ways, and hence one may argue the definition is still not precise enough. Today, cloud computing commonly refers to a computing model where services are hosted using resources in data centers and delivered to end users over the Internet. In our opinion, since cloud computing technologies are still evolving, finding the precise definition of cloud computing at the current moment may not be the right approach. Perhaps once the technologies have reached maturity, the true definition will naturally emerge.

1.2.2 Why Cloud Computing?

In this section, we present the motivation behind the development of cloud computing. We will also compare cloud computing with other parallel and distributed computing models and highlight their differences.

1.2.2.1 Key Driving Forces. There are several driving forces behind the success of cloud computing. The increasing demand for large-scale computation and big data analytics and economics are the most important ones. But other factors such as easy access to computation and storage, flexibility in resource allocations, and scalability play important roles.

Large-scale computation and big data: Recent years have witnessed the rise of Internet-scale applications. These applications range from social networks (e.g., facebook, twitter), video applications (e.g., Netflix, youtube), enterprise applications (e.g., Salesforce, Microsoft CRM) to personal applications (e.g., iCloud, Dropbox). These applications are commonly accessed by large numbers of users over the Internet. They are extremely large scale and resource intensive. Furthermore, they often have high performance requirements such as response time. Supporting these applications requires extremely large-scale infrastructures. For instance, Google has hundreds of compute clusters deployed worldwide with hundreds of thousands of servers. Another salient

characteristic is that these applications also require access to huge volumes of data. For instance, Facebook stores tens of petabytes of data and processes over a hundred terabytes per day. Scientific applications (e.g., brain image processing, astrophysics, ocean monitoring, and DNA analysis) are more and more deployed in the cloud. Cloud computing emerged in this context as a computing model designed for running large applications in a scalable and cost-efficient manner by harnessing massive resource capacities in data centers and by sharing the data center resources among applications in an on-demand fashion.

Economics: To support large-scale computation, cloud providers rely on inexpensive commodity hardware offering better scalability and performance/price ratio than supercomputers. By deploying a very large number of commodity machines, they leverage economies of scale bringing per unit cost down and allowing for incremental growth. On the other hand, cloud customers such as small and medium enterprises, which outsource their IT infrastructure to the cloud, avoid upfront infrastructure investment cost and instead benefit from a pay-as-you-go pricing and billing model. They can deploy their services in the cloud and make them quickly available to their own customers resulting in short time to market. They can start small and scale up and down their infrastructure based on their customers demand and pay based on usage.

Scalability: By harnessing huge computing and storage capabilities, cloud computing gives customers the illusion of infinite resources on demand. Customers can start small and scale up and down resources as needed.

Flexibility: Cloud computing is highly flexible. It allows customers to specify their resource requirements in terms of CPU cores, memory, storage, and networking capabilities. Customers are also offered the flexibility to customize the resources in terms of operating systems and possibly network stacks.

Easy access: Cloud resources are accessible from any device connected to the Internet. These devices can be traditional workstations and servers or less traditional devices such as smart phones, sensors, and appliances. Applications running in the cloud can be deployed or accessed from anywhere at anytime.

1.2.2.2 Relationship with Other Computing Models. Cloud computing is not a completely new concept and has many similarities with existing distributed and parallel computing models such as Grid computing and Cluster computing. But cloud computing also has some distinguishing properties that explain why existing models are not used and justify the need for a new one. These can be explained according to two dimensions: scale and service-orientation. Both parallel computing and cloud computing are used to solve large-scale problems often by subdividing these problems into smaller parts and carrying out the calculations concurrently on different processors. In the cloud, this is achieved using computational models such as MapReduce. However, while parallel computing relies on expensive supercomputers and massively parallel multi-processor machines, cloud computing uses cheap, easily replaceable commodity hardware. Grid computing uses supercomputers but can also use commodity hardware, all accessible through open, general-purpose protocols and interfaces, and distributed management and job scheduling middleware. Cloud computing differs from Grid computing in that

it provides high bandwidth between machines, that is more suitable for I/O-intensive applications such as log analysis, Web crawling, and big-data analytics. Cloud computing also differs from Grid computing in that resource management and job scheduling is centralized under a single administrative authority (cloud provider) and, unless this evolves differently in the future, provides no standard application programming interfaces (APIs). But perhaps the most distinguishing feature of cloud computing compared to previous computing models is its extensive reliance on virtualization technologies to allow for efficient sharing of resources while guaranteeing isolation between multiple cloud tenants. Regarding the second dimension, unlike other computing models designed for supporting applications and are mainly application-oriented, cloud computing extensively leverages service orientation providing everything (infrastructure, development platforms, software, and applications) as a service.

1.2.3 Architecture

Generally speaking, the architecture of a cloud computing environment can be divided into four layers: the hardware/datacenter layer, the infrastructure layer, the platform layer, and the application layer, as shown in Figure 1.1. We describe each of them in detail in the text that follows:

The hardware layer: This layer is responsible for managing the physical resources of the cloud, including physical servers, routers, and switches, and power, and cooling systems. In practice, the hardware layer is typically implemented in data centers. A data center usually contains thousands of servers that are organized in racks and interconnected through switches, routers, or other fabrics. Typical issues at hardware layer include hardware configuration, fault-tolerance, traffic management, and power and cooling resource management.

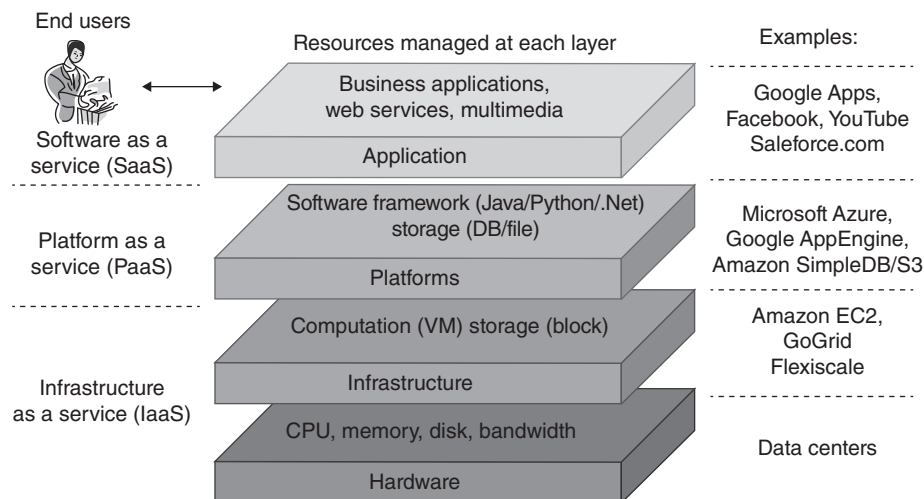


Figure 1.1. Typical architecture in a cloud computing environment.

The infrastructure layer: Also known as the virtualization layer, the infrastructure layer creates a pool of storage and computing resources by partitioning the physical resources using virtualization technologies such as Xen [4], KVM [5], and VMware [6]. The infrastructure layer is an essential component of cloud computing, since many key features, such as dynamic resource assignment, are only made available through virtualization technologies.

The platform layer: Built on top of the infrastructure layer, the platform layer consists of operating systems and application frameworks. The purpose of the platform layer is to minimize the burden of deploying applications directly into VM containers. For example, Google App Engine operates at the platform layer to provide API support for implementing storage, database, and business logic of typical Web applications.

The application layer: At the highest level of the hierarchy, the application layer consists of the actual cloud applications. Different from traditional applications, cloud applications can leverage the automatic-scaling feature to achieve better performance, availability, and lower operating cost. Compared to traditional service hosting environments such as dedicated server farms, the architecture of cloud computing is more modular. Each layer is loosely coupled with the layers above and below, allowing each layer to evolve separately. This is similar to the design of the protocol stack model for network protocols. The architectural modularity allows cloud computing to support a wide range of application requirements while reducing management and maintenance overhead.

1.2.4 Cloud Services

Cloud computing employs a service-driven business model. In other words, hardware and platform-level resources are provided as services on an on-demand basis. Conceptually, every layer of the architecture described in the previous section can be implemented as a service to the layer above. Conversely, every layer can be perceived as a customer of the layer below. However, in practice, clouds offer services that can be grouped into three categories: software as a service (SaaS), platform as a service (PaaS), and infrastructure as a service (IaaS).

1. *Infrastructure as a service:* IaaS refers to on-demand provisioning of infrastructural resources, usually in terms of VMs. The cloud owner who offers IaaS is called an IaaS provider.
2. *Platform as a service:* PaaS refers to providing platform layer resources, including operating system support and software development frameworks.
3. *Software as a service:* SaaS refers to providing on-demand applications over the Internet.

The business model of cloud computing is depicted in Figure 1.2. According to the layered architecture of cloud computing, it is entirely possible that a PaaS provider runs its cloud on top of an IaaS providers cloud. However, in the current practice, IaaS and

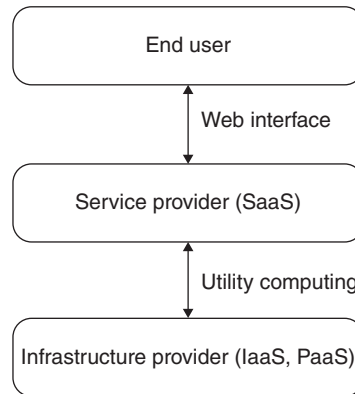


Figure 1.2. Cloud computing business model.

PaaS providers are often parts of the same organization (e.g., Google). This is why PaaS and IaaS providers are often called cloud providers [7].

1.2.4.1 Type of Clouds. There are many issues to consider when moving an enterprise application to the cloud environment. For example, some enterprises are mostly interested in lowering operation cost, while others may prefer high reliability and security. Accordingly, there are different types of clouds, each with its own benefits and drawbacks:

- *Public clouds:* A cloud in which cloud providers offer their resources as services to the general public. Public clouds offer several key benefits to service providers, including no initial capital investment on infrastructure and shifting of risks to cloud providers. However, current public cloud services still lack fine-grained control over data, network and security settings, which hampers their effectiveness in many business scenarios.
- *Private clouds:* Also known as internal clouds, private clouds are designed for exclusive use by a single organization. A private cloud may be built and managed by the organization or by external providers. A private cloud offers the highest degree of control over performance, reliability, and security. However, they are often criticized for being similar to traditional proprietary server farms and do not provide benefits such as no up-front capital costs.
- *Hybrid clouds:* A hybrid cloud is a combination of public and private cloud models that tries to address the limitations of each approach. In a hybrid cloud, part of the service infrastructure runs in private clouds while the remaining part runs in public clouds. Hybrid clouds offer more flexibility than both public and private clouds. Specifically, they provide tighter control and security over application data compared to public clouds, while still facilitating on-demand service expansion

and contraction. On the down side, designing a hybrid cloud requires carefully determining the best split between public and private cloud components.

- *Community clouds*: A community cloud refers to a cloud infrastructure that is shared between multiple organizations that have common interests or concerns. Community clouds are a specific type of cloud that relies on the common interest and limited participants to achieve efficient, reliable, and secure design of the cloud infrastructure.

Private cloud has always been the most popular type of cloud. Indeed, the development of cloud computing was largely due to the need of building data centers for hosting large-scale online services owned by large private companies, such as Amazon and Google. Subsequently, realizing the cloud infrastructure can be leased to other companies for profits, these companies have developed public cloud services. This development has also led to the creation of hybrid clouds and Community clouds, which represent different alternatives to share cloud resources among service providers. In the future, it is believed that private cloud will remain to be the dominant cloud computing model. This is because as online services continue to grow in scale and complexity, it becomes increasingly beneficial to build private cloud infrastructure to host these services. In this case, private clouds not only provide better performance and manageability than public clouds but also reduced operation cost. As the initial capital investment on a private cloud can be amortized across large number of machines over many years, in the long-term private cloud typically has lower operational cost compared to public clouds.

1.2.4.2 SME's Survey on Cloud Computing. The European Network and Information Security Agency (ENISA) has conducted a survey on the adaption of the cloud computing model by small to medium enterprises (SMEs). The survey provides an excellent overview of the benefits and limitations of today's cloud technologies. In particular, the survey has found that the main reason for adopting cloud computing is to reduce total capital expenditure on software and hardware resources. Furthermore, most of the enterprises prefer a mixture of cloud computing models (public cloud, private cloud), which comes with no surprise as each type of cloud has own benefits and limitations. Regarding the type of cloud services, it seems that IaaS, PaaS, and SaaS all received similar scores, even though SaaS is slightly in favor compared to the other two. Last, it seems that data availability, privacy, and confidentiality are the main concerns of all the surveyed enterprises. As a result, it is not surprising to see that most of the enterprises prefer to have a disaster recovery plan when considering migration to the cloud. Based on these observations, cloud providers should focus more on improving the security and reliability aspect of cloud infrastructures, as they represent the main obstacles for adopting the cloud computing model by today's enterprises.

1.2.5 Enabling Technologies

The success of cloud computing is largely driven by successful deployment of its enabling technologies. In this section, we provide an overview of cloud enabling technologies and describe how they contribute to the development of cloud computing.

1.2.5.1 Data Center Virtualization. One of the main characteristics of cloud computing is that the infrastructure (e.g., data centers) is often shared by multiple tenants (e.g., service providers) running applications with different resource requirements and performance objectives. This raises the question regarding how data center resources should be allocated and managed by each service provider. A naive solution that has been implemented in the early days is to allocate dedicated servers for each application. While this “bare-metal” strategy certainly worked in many scenarios, it also introduced many inefficiencies. In particular, if the server resource is not fully utilized by the application running on the server, the resource is wasted as no other application has the right to acquire the resource for its own execution. Motivated by this observation, the industry has adopted virtualization in today’s cloud data centers. Generally speaking, virtualization aims at partitioning physical resources into virtual resources that can be allocated to applications in a flexible manner. For instance, server virtualization is a technology that partitions the physical machine into multiple VMs, each capable of running applications just like a physical machine. By separating logical resources from the underlying physical resources, server virtualization enables flexible assignment of workloads to physical machines. This not only allows workload running on multiple VMs to be consolidated on a single physical machine, but also enables a technique called VM migration, which is the process of dynamically moving a VM from one physical machine to another. Today, virtualization technologies have been widely used by cloud providers such as Amazon EC2, Rackspace, and GoGrid. By consolidating workload using fewer machines, server virtualization can deliver higher resource utilization and lower energy consumption compared to allocating dedicated servers for each application.

Another type of data center virtualization that has been largely overlooked in the past is network virtualization. Cloud applications today are becoming increasingly data-intensive. As a result, there is a pressing need to determine how data center networks should be shared by multiple tenants with diverse performance, security and manageability requirements. Motivated by these limitations, there is an emerging trend towards virtualizing data center networks in addition to server virtualization. Similar to server virtualization, network virtualization aims at creating multiple VNs on top of a shared physical network substrate allowing each VN to be implemented and managed independently. By separating logical networks from the underlying physical network, it is possible to implement network resource guarantee and introduce customized network protocols, security, and management policies. Combining with server virtualization, a fully virtualized data centers support the allocation in the form of virtual infrastructures or VIs (also known as virtual data centers (VDC)), which consist of VMs inter-connected by virtual networks. The scheduling and management of VIs have been studied extensively in recent years. Commercial cloud providers are also pushing towards this direction. For example, the Amazon Virtual Private Cloud (VPC) already provides limited features to support network virtualization in addition to server virtualization.

1.2.5.2 Cloud Networking. To ensure predictable performance over the cloud, it is of utmost importance to design efficient networks that are able to provide guaranteed performance and to scale with the ever-growing traffic volumes in the cloud. Traditional

data center network architectures suffer from many limitations that may hinder the performance of large-scale cloud services. For instance, the widely-used tree-like topology does not provide multiple paths between the nodes, and hence limits the scalability of the network and the ability to mitigate node and link congestion and failures. Moreover, current technologies like Ethernet and VLANs are not well suited to support cloud computing requirements like multi-tenancy or performance isolation between different tenants/applications. In recent years, several research works have focused on designing new data center network architectures to overcome these limitations and enhance performance, fault tolerance and scalability (e.g., VL2 [38], Portland [9], NetLord [10]). Furthermore, the advent of software-defined networking (SDN) technology brings new opportunities to redesign cloud networks [11]. Thanks to the programmability offered by this technology, it is now possible to dynamically adapt the configuration of the network based on the workload. It also makes it easy to implement policy-based network management schemes in order to achieve potential cloud providers’ objectives in terms of performance, utilization, survivability, and energy efficiency.

1.2.5.3 Data Storage and Management. As mentioned previously, one of the key driving forces for cloud computing is the need to process large volumes of data in a scalable and efficient manner. As cloud data centers typically consist of commodity servers with limited storage and processing capacities, it is necessary to develop distributed storage systems that support efficient retrieval of desired data. At the same time, as failures are common in commodity machine-based data centers, the distributed storage system must also be resilient to failures. This usually implies each file block must be replicated on multiple machines. This raises challenges regarding how the distributed storage system should be designed to achieve availability and high performance, while ensuring file replicas remain consistent over time. Unfortunately, the famous CAP theorem [12] states that simultaneously achieving all three objectives (consistency, availability, and robustness to network failures) is not a viable task. As result, recently many file systems such Google File System [13], Amazon Dynamo [14], Cassandra [15] are trying to explore various trade-offs among the three objectives based on applications’ needs. For example, Amazon Dynamo adopts an eventual consistency model that allow replicas to be temporary out-of-sync. By sacrificing consistency, Dynamo is able to achieve significant improvement in server response time. It is evident that these storage systems provide the foundations for building large-scale data-intensive applications that are commonly found in today’s cloud data centers.

1.2.5.4 MapReduce Programming Model. Cloud computing has become the most cost-effective technology for hosting Internet-scale applications. Companies like Google and Facebook generate enormous volumes of data on a daily basis that need to be processed in a timely manner. To meet this requirement, cloud providers use computational models such as MapReduce [1] and Dryad [16]. In these models, a job spawns many small tasks that can be executed concurrently on multiple machines, resulting in significant reduction in job completion time. Furthermore, to cope with software and hardware exceptions frequent in large-scale clusters, these models provide built-in fault

tolerance features that automatically restart failed tasks when exceptions occur. As a result, these computational models are very attractive not only for running data-intensive jobs but also for computation-intensive applications. The MapReduce model, in particular, is largely used nowadays in cloud infrastructures for supporting a wide range of applications and has been adapted to several computing and cluster environments. Despite this success, the adoption of MapReduce has implications on the management of cloud workload and cluster resources, which is still largely unstudied. In particular, many challenges pertaining to MapReduce job scheduling, task and data placement, resource allocation, and sharing are yet to be addressed.

1.2.5.5 Resource Management. Resource management has always been a central theme of cloud computing. Given the large variety of applications running in the cloud, it is a challenging problem to determine how each application should be scheduled and managed in a scalable and dynamic manner. The scheduling of individual application component can be formulated as a variant of the multi-dimensional vector bin-packing problem, which is already NP-hard in the general case. Furthermore, different applications may have different scheduling needs. For example, individual tasks of a single MapReduce job can be scheduled independently over time, whereas the servers of a three-tier Web application must be scheduled simultaneously to ensure service availability. Therefore, finding a scheduling scheme that satisfy diverse application scheduling requirement is a challenging problem. The recent work on multi-framework scheduling (e.g., MESOS [17]) provides a platform to allow various scheduling frameworks, such as MapReduce, Spark, and MPI to coexist in a single cloud infrastructure. The work on distributed schedulers (e.g., Omega [18] and Sparrow [19]) also aim at improving the scalability of schedulers by having multiple schedulers perform scheduling in parallel. These technologies will provide the functionality to support a wide range of workload in the cloud data center environments.

1.2.5.6 Energy Management. Data centers consume tremendous amount of energy, not only for powering up the servers and network devices, but also for cooling down these components to prevent overheating conditions. It has been reported that energy cost accounts for 15% of the average data center operation expenditure. At the same time, such large energy consumption also raises environmental concerns regarding the carbon emissions for energy generation. As a result, improving data center energy efficiency has become a primary concern for today’s data center operators. A widely used metric for measuring energy efficiency of data centers is power usage effectiveness (PUE), which is computed as the ratio between the computer infrastructure usage and the total data center power usage. Even though none of the existing data centers can achieve the ideal PUE value of 1.0, many cloud data centers today have become very energy efficient with PUE less than 1.1.

There are many techniques for improving data center energy efficiency. At the infrastructure level, many cloud providers leverage nearby renewable energy source (i.e., solar and wind) to reduce energy cost and carbon footprint. At the same time, it is also possible to leverage environmental conditions (e.g., low temperature conditions) to reduce

cooling cost. For example, Facebook recently announced the construction of a cloud data center in Sweden, right on the edge of the arctic circle, mainly due to the low air temperature that can reduce cooling cost. The Net-Zero Energy Data Center developed by HP labs leverages locally generated renewable energy and workload demand management techniques to significantly reduce the energy required to operate data centers. We believe the rapid development of cloud energy management techniques will continue to push the data center energy efficiency towards the ideal PUE value of 1.0.

1.2.5.7 Security and Privacy. Security is another major concern of cloud computing. While security is not a critical concern in many private clouds, it is often a key barrier to the adoption of cloud computing in public clouds. Specifically, since service providers typically do not have access to the physical security system of data centers, they must rely on cloud providers to achieve full data security. The cloud provider, in this context, must achieve the following objectives: (1) confidentiality, for secure data access and transfer, and (2) auditability, for attesting whether security setting of applications has been tampered or not. Confidentiality is usually achieved using cryptographic protocols, whereas auditability can be achieved using remote attestation techniques. Remote attestation typically requires a trusted platform module (TPM) to generate nonforgeable system summary (i.e., system state encrypted using TPM private key) as the proof of system security. However, in a virtualized environment like the clouds, VMs can dynamically migrate from one location to another, hence directly using remote attestation is not sufficient. In this case, it is critical to build trust mechanisms at every architectural layer of the cloud. First, the hardware layer must be trusted using hardware TPM. Second, the virtualization platform must be trusted using secure VM monitors. VM migration should only be allowed if both source and destination servers are trusted. Recent work has been devoted to designing efficient protocols for trust establishment and management.

1.3 PART II: RESEARCH CHALLENGES—THE CHAPTERS IN THIS BOOK

This book covers the fundamentals of cloud services, networking and management and focuses on most prominent research challenges that have drawn the attention of the IT community in the past few years. Each of the 14 chapters of this book provides an overview of some of the key architectures, features, and technologies of cloud services, networking and management systems and highlights state-of-the-art solutions and possible research gaps. The chapters of the book are written by knowledgeable authors that were carefully selected based on their expertise in the field. Each chapter went through a rigorous review process, including external reviewers, the book editors Raouf Boutaba and Nelson Fonseca, and the series editors Tom Plevyak and Veli Sahin. In the following, we briefly describe the topics covered by the different chapters of this book.

1.3.1 Virtualization in the Cloud

Virtualization is one of the key enabling technologies that made cloud computing model a reality. Initially, virtualization technologies have allowed to partition a physical server

into multiple isolated environments called VMs that may eventually host different operating systems and be used by different users or applications. As cloud computing evolved, virtualization technologies have matured and have been extended to consider not only the partitioning of servers but also the partitioning of the networking resources (e.g., links, switches and routers). Hence, it is now possible to provide each cloud user with a VI encompassing VMs, virtual links, and virtual routers and switches. In this context, several challenges arise especially regarding the management of the resulting virtualized environment where different types of resources are shared among multiple users.

In this chapter, the authors outline the main characteristics of these virtualized infrastructures and shed light on the different management operations that need to be implemented in such environments. They then summarize the ongoing efforts towards defining open standard interfaces to support virtualization and interoperability in the cloud. Finally, the chapter provides a brief overview of the main open-source cloud management platforms that have recently emerged.

1.3.2 VM Migration

One of the powerful features brought by virtualization is the ability to easily migrate VMs within the same data center or even between geographically distributed data centers. This feature provides an unprecedented flexibility to network and data center operators allowing them to perform several management tasks like dynamically optimizing resource allocations, improving fault tolerance, consolidating workloads, avoiding server overload, and scheduling maintenance activities. Despite all these benefits, VM migration induces several costs, including higher utilization of computing and networking resources, inevitable service downtime, security risks, and more complex management challenges. As a result, a large number of migration techniques have been recently proposed in the literature in order to minimize these costs and make VM migration a more effective and secure tool in the hand of cloud providers.

This chapter starts by providing an overview of VM migration techniques. It then presents, XenFlow, a tool based on Xen and OpenFlow, and allowing to deploy, isolate and migrate VIs. Finally, the authors discuss potential security threats that can arise when using VM migration.

1.3.3 Data Center Networks and Relevant Standards

Today’s cloud data centers are housing hundreds of thousands of machines that continuously need to exchange tremendous amounts of data with stringent performance requirements in terms of bandwidth, delay, jitter, and loss rate. In this context, the data center network plays a central role to ensure a reliable and efficient communication between machines, and thereby guarantee continuous operation of the data center and effective delivery of the cloud services. A data center network architecture is typically defined by the network topology (i.e., the way equipment are inter-connected) as well as the adopted switching, routing, and addressing schemes and protocols (e.g., Ethernet and IP).

Traditional data center network architectures suffer from several limitations and are not able to satisfy new application requirements spawned by cloud computing model in terms of scalability, multitenancy and performance isolation. For instance, the widely used tree-like topology does not provide multiple paths between the nodes, and hence limits the ability to survive node and link failures. Also, current switches have limited forwarding table sizes, making it difficult for traditional data center networks to handle the large number of VMs that may exist in virtualized cloud environments. Another issue is with the performance isolation between tenants as there is no bandwidth allocation mechanism in place to ensure predictable network performance for each of them.

In order to cope with these limitations, a lot of attention has been devoted in the past few years to study the performance of existing architectures and to design better solutions. This chapter dwells on these solutions covering data center network architectures, topologies, routing protocols and addressing schemes that have been recently proposed in the literature.

1.3.4 Interdata Center Networks

In recent years, cloud providers have largely relied on large-scale cloud infrastructures to support Internet-scale applications efficiently. Typically, these infrastructures are composed of several geographically distributed data centers connected through a backbone network (i.e., an inter-data center network). In this context, a key challenge facing cloud providers is to build cost-effective backbone networks while taking into account several considerations and requirements including scalability, energy efficiency, resilience, and reliability. To address this challenge, many factors should be considered. The scalability requirement is due to the fact that the volume of data exchanged between data centers is growing exponentially with the ever-increasing demand in cloud environments. The energy efficiency requirement concerns how to minimize the energy consumption of the infrastructure. Such a requirement is not only crucial to make the infrastructure more green and environmental-friendly but also essential to cut down operational expenses. Finally, the resilience of the interdata center network requirement is fundamental to maintain a continuous and reliable cloud services.

This chapter investigates the different possible alternatives to design and manage cost-efficient cloud backbones. It then presents mathematical formulations and heuristic solutions that could be adopted to achieve desired objectives in terms of energy efficiency, resilience and reliability. Finally, the authors discuss open issues and key research directions related to this topic.

1.3.5 OpenFlow and SDN for Clouds

The past few years have witnessed the rise of SDN, a technology that makes it possible to dynamically configure and program networking elements. Combined with cloud computing technologies, SDN enables the design of highly dynamic, efficient, and cost-effective shared application platforms that can support the rapid deployment of Internet applications and services.

This chapter discusses the challenges faced to integrate SDN technology in cloud application platforms. It first provides a brief overview of the fundamental concepts of SDN including OpenFlow technology and tools like Open vSwitch. It also introduces the cloud platform OpenStack with a focus on its Networking Service (i.e., Neutron project), and shows how cloud computing environments can benefit from SDN technology to provide guaranteed networking resources within a data center and to interconnect data centers. The authors also review major open source efforts that attempt to integrate SDN technology in cloud management platforms (e.g., OpenDaylight open source project) and discuss the notion of software-defined infrastructure (SDI).

1.3.6 Mobile Cloud Computing

Mobile cloud computing has recently emerged as a new paradigm that combines cloud computing with mobile network technology with the goal of putting the scalability and limitless resources of the cloud into the hands of mobile service and application providers. However, despite of its potential benefits, the growth of mobile cloud computing in recent years was hampered by several technical challenges and risks. These challenges and risks are mainly due to the inherent limitations of mobile devices such as the scarcity of resources, the limited energy supply, the intermittent connectivity in wireless networks, security risks, and legal/environmental risks.

This chapter starts by providing an overview of mobile cloud computing application models and frameworks. It also defines risk management and identifies and analyzes prevalent risk factors found in mobile cloud computing environments. The authors also present an analysis of mobile cloud frameworks from a risk management perspective and discusses the effectiveness of traditional risk approaches to address mobile cloud computing risks.

1.3.7 Resource Management and Scheduling

Resource allocation and scheduling are two crucial functions in cloud computing environments. Generally speaking, cloud providers are responsible for allocating resources (e.g., VMs) with the goal of satisfying the promised service-level agreement (SLA) while increasing their profit. This can be achieved by reducing operational costs (e.g., energy costs) and sharing resources among the different users. At the opposite side, cloud users are responsible for application scheduling that aims at mapping tasks from applications submitted by users to computational resources in the system. The goals of scheduling include maximizing the usage of the leased resources, and minimizing costs by dynamically adjusting the leased resources to the demand while maintaining the required quality of service.

Resource allocation and scheduling are both vital to cloud users and providers, but they both have their own specifics, challenges and potentially conflicting objectives. This chapter starts by a review of the different cloud types and service models and then discusses the typical objectives of cloud providers and their clients. The chapter provides also mathematical formulations to the problems, VM allocation, and application

scheduling. It surveys some of the existing solutions and discusses their strengths and weaknesses. Finally, it points out the key research directions pertaining to resource management in cloud environments.

1.3.8 Autonomic Performance Management for Multi-Clouds

The growing popularity of the cloud computing model have led to the emergence of multiclouds or clouds of clouds where multiple cloud systems are federated together to further improve and enhance cloud services. Multiclouds have several benefits that range from improving availability, to reducing lock-in, and optimizing costs beyond what can be achieved within a single cloud. At the same time, multi-clouds bring new challenges in terms of the design, development, deployment, monitoring, and management of multi-tier applications able to capitalize on the advantages of such distributed infrastructures. As a matter of fact, the responsibility for addressing these challenges is shared among cloud providers and cloud users depending on the type of service (i.e., IaaS, PaaS, and SaaS) and SLAs. For instance, from an IaaS cloud provider’s perspective, management focuses mainly on maintaining the infrastructure, allocating resources requested by clients and ensuring their high availability. By contrast, cloud users are responsible for implementing, deploying and monitoring applications running on top of resources that are eventually leased from several providers. In this context, a compelling challenge that is currently attracting a lot of attention is how to develop sophisticated tools that simplify the process of deploying, managing, monitoring, and maintaining large-scale applications over multi-clouds.

This chapter focuses on this particular challenge and provides a detailed overview of the design and implementation of XCAMP, the X-Cloud Application Management Platform that allows to automate application deployment and management in multitier clouds. It also highlights key research challenges that require further investigation in the context of performance management and monitoring in distributed cloud environments.

1.3.9 Energy Management

Cloud computing environments mainly consist of data centers where thousands of servers and other systems (e.g. power distribution and cooling equipment) are consuming tremendous amounts of energy. Recent reports have revealed that energy costs represent more than 12% of the total data center operational expenditures, which translates into millions of dollars. More importantly, high energy consumption is usually synonymous of high carbon footprint, raising serious environmental concerns and pushing governments to put in place more stringent regulations to protect the environment. Consequently, reducing energy consumption has become one of the key challenges facing today’s data center managers. Recently, a large body of work has been dedicated to investigate possible techniques to achieve more energy-efficient and environment-friendly infrastructures. Many solutions have been proposed including dynamic capacity provisioning and optimal usage of renewable sources of energy (e.g., wind power and solar).

This chapter further details the trends in energy management solutions in cloud data centers. It first surveys energy-aware resource scheduling and allocation schemes aiming at improving energy efficiency, and then provides a detailed description of GreenCloud, an energy-aware cloud data center simulator.

1.3.10 Survivability and Fault Tolerance in the Cloud

Despite the success of cloud computing, its widespread and full-scale adoption have been hampered by the lack of strict guarantees on the reliability and availability of the offered resources and services. Indeed, outages, failures and service disruption can be fatal for many businesses. Not only they incur significant revenue loss—as much as hundreds of thousands of dollars per minute for some services—but they may also hurt the business reputation in the long term and impact on customers’ loyalty and satisfaction. Unfortunately, major cloud providers like Amazon EC2, Google, and Rackspace are not yet able to satisfy the high availability and reliability levels required for such critical services.

Consequently, a growing body of work has attempted to address this problem and to propose solutions to improve the reliability of cloud services and eventually provide more stringent guarantees to cloud users. This chapter provides a comprehensive literature survey on this particular topic. It first lays out cloud computing and survivability-related concepts, and then covers recent studies that analyzed and characterized the types of failures found in cloud environments. Subsequently, the authors survey and compare the solutions put forward to enhance cloud services’ fault-tolerance and to guarantee high availability of cloud resources.

1.3.11 Cloud Security

Security has always been a key issue for cloud-based services and several solutions have been proposed to protect the cloud from malicious attacks. In particular, intrusion detection systems (IDS) and intrusion prevention systems (IPS) have been widely deployed to improve cloud security and have been recently empowered with new technologies like SDN to further enhance their effectiveness. For instance, the SDN technology has been leveraged to dynamically reconfigure the cloud network and services and better protect them from malicious traffic. In this context, this chapter introduces FlowIPS, an OpenFlow-based IPS solution for intrusion prevention in cloud environments. FlowIPS implements SDN-based control functions based on Open vSwitch (OVS) and provides novel Network Reconfiguration (NR) features by programming POX controllers. Finally, the chapter presents the performance evaluation of FlowIPS that demonstrates its efficiency compared to traditional IPS solutions.

1.3.12 Big Data on Clouds

Big data has emerged as a new term that describes all challenges related to the manipulation of large amounts of data including data collection, storage, processing, analysis, and visualization.

This chapter articulates some of the success enablers for deploying Big Data on Clouds (BDOC). It starts by providing some historical perspectives and by describing emerging Internet services and applications. It then describes some legal issues related to big data on clouds. In particular, it highlights emerging hybrid big data management roles, the development and operations (DevOps), and Site Reliability Engineering (SRE). Finally, the chapter discusses science, technology, engineering, and mathematics (STEM) talent cultivation and engagement, as an enabler to technical succession and future success for global enterprises of big data on clouds.

1.3.13 Scientific Applications on Clouds

In order to cope with the requirements of scientific applications, cloud providers have recently proposed new coordination and management tools and services (e.g., Amazon Simple WorkFlow or SWF) in order to automate and streamline task processes executed by the cloud applications. Such services allow to specify the dependencies between the tasks, their order of execution and make it possible to track their progress and the current state of each of them. In this context, a compelling challenge is to ensure the compatibility between existing workflow systems and to provide the possibility to reuse scientific legacy code.

This chapter presents a software engineering solution that allows the scientific workflow community to use Amazon cloud via a single front-end converter. In particular, it describes a wrapper service for executing legacy code using Amazon SWF. The chapter also describes the experimental results demonstrating that the automatically SWF application generated by the wrapper provides a performance comparable to the native manually optimized workflow.

1.3.14 Interactive Multimedia Applications on Clouds

The booming popularity of cloud computing has led to the emergence of a large array of new applications such as social networking, gaming, live streaming, TV broadcasting, and content delivery. For instance, cloud gaming allows direct on-demand access to games whose content is stored in the cloud and streamed directly to end users through thin clients. As a result, less powerful game consoles or computers are needed as most of the processing is carried out in the hosting cloud, leveraging its seemingly unlimited resources. Another prominent cloud application is the Massive user-generated content (UGC) live streaming that allows each simple Internet user to become a TV or content provider. A similar application that has become extremely popular is time-shifting on-demand TV as many services like catch-up TV (i.e., the content of a TV channel is recorded for many days and can be requested on demand) and TV surfing (i.e., the possibility of pausing, forwarding or rewinding of a video stream) have recently become widely demanded. Naturally, the cloud is the ideal platform to host such services as it provides the processing and storage capacity required to ensure a high quality of service. However, several challenges are not addressed yet especially because of the stringent performance requirements (e.g., delay) of such multimedia applications and the increasing amounts of traffic they generate.

REFERENCES

21

This chapter discusses the deployment of these applications over the cloud. It starts by laying out content delivery models in general, and then provides a detailed study of the performance of three prominent multimedia cloud applications, namely cloud gaming, massive user-generated content live streaming and time-shifting on-Demand TV.

1.4 CONCLUSION

Editing and preparing a book on such an important topic is a challenging task requiring a lot of effort and time. As the editors of this book, we are grateful to many individuals who contributed to its successful completion. We would like to thank the chapters’ authors for their high-quality contributions, the reviewers for their insightful comments and feedback, and the book series editors for their support and guidance. Finally, we hope that the reader finds the topics and the discussions presented in this book informative, interesting, and inspiring and pave the way for designing new cloud platforms able to meet the requirements of future Internet applications and services.

REFERENCES

1. J. Dean and S. Ghemawat, “MapReduce: Simplified data processing on large clusters,” *Communications of the ACM*, vol. 51, no. 1, pp. 107–113, 2008.
2. L. M. Vaquero, L. Roderio-Merino, J. Caceres, and M. Lindner, “A break in the clouds: Towards a cloud definition,” *ACM SIGCOMM Computer Communication Review*, vol. 39, no. 1, pp. 50–55, 2008.
3. P. Mell and T. Grance, “The NIST definition of cloud computing (draft),” *NIST Special Publication*, vol. 800, no. 145, p. 7, 2011.
4. P. Barham, B. Dragovic, K. Fraser, S. Hand, T. Harris, A. Ho, R. Neugebauer, I. Pratt, and A. Warfield, “Xen and the art of virtualization,” *ACM SIGOPS Operating Systems Review*, vol. 37, no. 5, pp. 164–177, 2003.
5. A. Kivity, Y. Kamay, D. Laor, U. Lublin, and A. Liguori, “KVM: The linux virtual machine monitor,” in *Proceedings of the Linux Symposium*, vol. 1, Dttawa, Dntorio, Canada, 2007, pp. 225–230.
6. F. Guthrie, S. Lowe, and K. Coleman, *VMware vSphere Design*. John Wiley & Sons, Indianapolis, IN, 2013.
7. A. Fox, R. Griffith, A. Joseph, R. Katz, A. Konwinski, G. Lee, D. Patterson, A. Rabkin, and I. Stoica, “Above the clouds: A berkeley view of cloud computing,” *Department of Electrical Engineering and Computer Sciences, University of California, Berkeley, CA, Rep. UCB/EECS*, vol. 28, p. 13, 2009.
8. A. Greenberg, J. Hamilton, N. Jain, S. Kandula, C. Kim, P. Lahiri, D. Maltz, P. Patel, and S. Sengupta, “VL2: A scalable and flexible data center network,” in *Proceedings ACM SIGCOMM*, Barcelona, Spain, August 2009.
9. R. Mysore, A. Pamboris, N. Farrington, N. Huang, P. Miri, S. Radhakrishnan, V. Subramanya, and A. Vahdat, “PortLand: A scalable fault-tolerant layer 2 data center network fabric,” in *Proceedings ACM SIGCOMM*, Barcelona, Spain, August 2009.

10. J. Mudigonda, P. Yalagandula, B. Stiekes, and Y. Pouffary, “NetLord: A scalable multi-tenant network architecture for virtualized datacenters,” in *Proceedings ACM SIGCOMM*, Toronto, Ontario, Canada, August 2011.
11. N. McKeown, T. Anderson, H. Balakrishnan, G. Parulkar, L. Peterson, J. Rexford, S. Shenker, and J. Turner, “Openflow: Enabling innovation in campus networks,” *SIGCOMM Computer Communication Review*, vol. 38, no. 2, pp. 69–74, March 2008.
12. S. Gilbert and N. Lynch, “Brewer’s conjecture and the feasibility of consistent, available, partition-tolerant web services,” *ACM SIGACT News*, vol. 33, no. 2, pp. 51–59, 2002.
13. S. Ghemawat, H. Gobioff, and S.-T. Leung, “The Google file system,” in *ACM SIGOPS Operating Systems Review*, vol. 37, no. 5, pp. 29–43, 2003.
14. G. DeCandia, D. Hastorun, M. Jampani, G. Kakulapati, A. Lakshman, A. Pilchin, S. Sivasubramanian, P. Voshall, and W. Vogels, “Dynamo: Amazon’s highly available key-value store,” in *ACM SIGOPS Operating Systems Review*, vol. 41, no. 6, pp. 205–220, 2007.
15. A. Lakshman and P. Malik, “Cassandra: A decentralized structured storage system,” *ACM SIGOPS Operating Systems Review*, vol. 44, no. 2, pp. 35–40, 2010.
16. M. Isard, M. Budiu, Y. Yu, A. Birrell, and D. Fetterly, “Dryad: Distributed data-parallel programs from sequential building blocks,” *ACM SIGOPS Operating Systems Review*, vol. 41, no. 3, pp. 59–72, 2007.
17. B. Hindman, A. Konwinski, M. Zaharia, A. Ghodsi, A. D. Joseph, R. Katz, S. Shenker, and I. Stoica, “MESOS: A platform for fine-grained resource sharing in the data center,” in *Proceedings of the 8th USENIX Conference on Networked Systems Design and Implementation*, Boston, MA, 2011, pp. 22–22.
18. M. Schwarzkopf, A. Konwinski, M. Abd-El-Malek, and J. Wilkes, “Omega: Flexible, scalable schedulers for large compute clusters,” in *Proceedings of the 8th ACM European Conference on Computer Systems*, Prague, Czech Republic. ACM, New York, 2013, pp. 351–364.
19. K. Ousterhout, P. Wendell, M. Zaharia, and I. Stoica, “Sparrow: Distributed, low latency scheduling,” in *Proceedings of the Twenty-Fourth ACM Symposium on Operating Systems Principles*, Farmington, PA. ACM, New York, 2013, pp. 69–84.