

RESEARCH ARTICLE

CloudSimSDN-NFV: Modeling and Simulation of Network Function Virtualization and Service Function Chaining in Edge Computing Environments

Jungmin Son | TianZhang He | Rajkumar Buyya*

Cloud Computing and Distributed Systems (CLOUDS) Laboratory, School of Computing and Information Systems, The University of Melbourne, VIC 3010, Australia

Correspondence

*Rajkumar Buyya, School of Computing and Information Systems, The University of Melbourne, VIC 3010, Australia. Email: rbuyya@unimelb.edu.au

Abstract

Software-defined networking (SDN) has evolved and brought an innovative paradigm shift in computer networks by utilizing a programmable software controller with open protocols. Network functions, previously served on dedicated hardware, have shifted to network function virtualization (NFV) which enabled functions to be virtualized and provisioned dynamically upon generic hardware. In addition to NFV, edge computing utilizes the edge resources close to end-users, which can reduce the end-to-end service delay and the network traffic volume. Although these innovative technologies gained significant attention from both academia and industry, there are limited tools and simulation frameworks for the effectiveness evaluation in a repeatable and controllable manner. Furthermore, large-scale experimental infrastructures are expensive to setup and difficult to maintain. Even if they are created, they are not available or accessible for the majority of researchers throughout the world. In this paper, we propose a framework for simulating NFV functionalities in both edge and cloud computing environments. In addition to the basic network functionalities supported by SDN in CloudSimSDN, we added new NFV features, such as virtual network functions (VNFs) allocation, migration, and auto-scaling with the support of corresponding network functionalities, such as flow load balancing, rerouting, and service function chaining (SFC) maintenance. We evaluated our simulation framework with auto-scaling and placement policies for SFC in the integrated edge and cloud computing environments. The results demonstrate its effectiveness in measuring and evaluating the end-to-end delay, response time, resource utilization, network traffic, and power consumption with different algorithms in each scenario.

KEYWORDS:

Cloud computing; Software-defined networking (SDN); Software-defined clouds; Network Function Virtualization (NFV); Service Function Chaining (SFC); Edge computing; Simulation software;

This is the author manuscript accepted for publication and has undergone full peer review but has not been through the copyediting, typesetting, pagination and proofreading process, which may lead to differences between this version and the Version of Record. Please cite this article as doi: [10.1002/spe.2755](https://doi.org/10.1002/spe.2755)

1 | INTRODUCTION

In the past decade, cloud computing has been evolved to elastically provision computing resources to their tenants in pay-as-you-go basis. Organizations and start-up companies can build their application services upon cloud data centers without investing huge up-front costs to purchase computing servers and network infrastructure. Instead, they can utilize as many resources as needed within minutes by several mouse-clicks from cloud service providers and pay only for the duration and the amount of provisioned resources¹. In cloud computing, physical machines are virtualized through hypervisors where the virtualized resources (virtual machines) can be allocated directly to the tenants or through platforms and services.

Similarly, virtualization technology has brought a new concept of network virtualization in the telecommunication field along with software-defined networking (SDN)². Network functions, such as firewall, proxy, and intrusion detection system (IDS), used to be served by an expensive hardware purpose-built only for certain network functions. As network functions are CPU intensive tasks, the network providers have to purchase the dedicated device to provide the required functions to their customers. Recently, network function virtualization (NFV)³ has been evolved drastically thanks to the advancement of virtualization technology, which we have seen in cloud computing. As virtual machines can be dynamically scaled, provisioned, and migrated in clouds, virtualized network functions (VNF) can be also provisioned throughout generic physical machines to provide a certain network function. Telecommunication providers build their own cloud data centers within their networks for NFV and place VNFs in their data center to elastically manage its computing and networking resources.

In addition to cloud computing, the increasing popularity of Internet-of-Things (IoT) leads to introducing edge and fog computing which utilizes more edge resources closer to the end users and IoT devices⁴. In IoT, the massive amount of network traffic generated by IoT sensors becomes challenging to service providers and network operators. In order to reduce the network traffic between the IoT sensor devices and the central computation servers usually residing in a cloud, the edge resources (e.g., network routers, wireless access points, and radio towers) can be utilized for filtering and pre-processing the collected data. Also, placing service functions nearby the end-users can reduce the end-to-end latency by reducing the network transmission delay between the end-user and the service function. In NFV, utilizing edge resources for VNFs has been studied by many researchers in recent years^{5,6,7,8}.

Despite the increased attention of NFV and edge computing, there are limited tools that can verify and evaluate new techniques in the literature. Several proof-of-concept platforms have been proposed by researchers^{9,10,8} to show the potential performance and orchestration abilities of NFV over edge computing. However, evaluating a new method in a large-scale is troublesome with the proposed platforms, because the empirical system has to be deployed to a large-scale infrastructure in order to evaluate at such scale. Traditional network simulation tools such as *ns3* does not support the new paradigm of NFV and SDN technologies. Thus, an accessible, adaptable, and scalable simulation toolkit have to be introduced and developed in order to foster emerging research and realization of NFV and edge computing in SDN-enabled cloud computing environments.

In order to fill the gap, in this paper, we propose a new simulation framework, **CloudSimSDN-NFV**, for NFV and edge computing simulation in SDN-enabled clouds extended from CloudSimSDN¹¹. We developed and presented CloudSimSDN in 2015 to simulate software-defined networking (SDN) functionalities in cloud computing. CloudSimSDN has helped to simulate different allocation and provisioning policies for both computing and networking resources¹², as well as the workload and application-aware scheduling methods in clouds¹³. In addition to the basic networking functionality supported by SDN in CloudSimSDN, this paper presents the design and implementation of new features to support NFV functionalities, such as VNF allocation, migration, and auto-scaling. To facilitate these functionalities, the simulation framework is based on the mapping of architecture and components of ETSI NFV NFV Management and orchestration (MANO)¹⁴, which includes NFV Orchestrator (NFVO), VNF Manager (VNFM), and Virtual Infrastructure Manager (VIM). Also, the concept of edge computing is integrated and supported by the new simulation framework, including multiple data centers, inter-cloud networks, and differentiated data center capacities.

The key **contributions** of this paper are:

- modeling of resource provisioning for NFV in the edge computing environment;

- architecture and design of the simulation framework for NFV in edge and cloud computing;
- detailed development and implementation of the simulation framework and the challenges;
- performance evaluation of the framework with use case scenarios;
- potential extensions and directions of the proposed framework.

The rest of the paper is organized as follows. Section 2 discusses existing platforms and simulation tools in the literature. Section 3 presents the modeling and simulation of NFV in edge and cloud computing environments, followed by detailed design and implementation of the new simulation framework in Section 4. Use case scenarios and evaluation results using the simulation framework are presented in Section 5. In Section 6, we discuss the potential extensions of proposed framework which can be implemented for supporting in different scenarios. Finally, Section 7 summarizes and concludes the paper .

2 | RELATED WORK

Several works have been proposed and presented in the literature to simulate cloud, edge, and fog computing, and networking. Many proof-of-concept systems are also developed for NFV evaluation. In this section, we review some of the related works and compare with our proposed framework.

Calheiros et al. developed CloudSim toolkit in 2011¹⁵ to simulate events and interactions of cloud data centers, such as virtual machine placement, provisioning resources, and scheduling workloads. It is a discrete event simulation tool, where every interaction is modeled as an event between cloud entities (e.g., cloud data center, broker, etc.). Simulation results are calculated based on the sending and receiving time of an event. Its simplicity and ease of use have attracted significant attention from both academia and industry, which led to initiating various descendant projects based on CloudSim, such as NetworkCloudSim¹⁶, ContainerCloudSim¹⁷, iFogSim¹⁸, CloudSim Plus¹⁹ and CloudSimSDN¹¹.

iFogSim¹⁸ was developed for simulation of fog computing environment. The concept of fog is similar to edge computing which utilizes not only central clouds but also edge resources closer to the end-users, whereas fog computing includes central cloud and other intermediate nodes in the architecture. With iFogSim, it is possible to create an integrated edge-cloud environment to evaluate resource management policies for both edge and clouds. However, it focuses on managing computing resources (CPU, memory, and storage) which result in limited functionality in networks such as NFV and dynamic network configuration. Our proposed framework, on the other hand, can simulate sophisticated network functionalities available in the SDN and NFV paradigm.

CloudSimSDN¹¹ was introduced in 2015 to provide simulation framework for SDN functionalities on the cloud computing environment. After the first presentation, it is utilized for evaluations in multiple projects, such as over-subscription based VM and network allocation policy¹², and priority-aware VM allocation policy considering network topology¹³. The tool can simulate dynamic network flow scheduling, joint VM and network optimization, and monitoring CPU and network utilization in both physical host level and virtual machine level. Although it provides extensive evaluation of network functionality with potential expandability, it is lack of supporting NFV and edge computing models.

Mininet²⁰ was developed by Stanford University to emulate SDN controllers in a single Linux machine. In Mininet, network switches and hosts are virtually created within the Linux operating system with the network virtualization functions provided by the Linux kernel. With Mininet, any OpenFlow-compatible SDN controllers can be used for emulating SDN functions with real-world traffic and scenarios. As it uses network drivers in the Linux kernel, the emulation reflects more accurate empirical results from the implemented SDN control logic. However, its scalability is limited due to the limitation of the operating system and its resource usage, which prevents specifically cloud-scale (e.g., thousands of machines) simulation.

Many researchers developed proof-of-concept systems to evaluate their approaches in NFV. LightMANO is proposed by Riggio et al.²¹ to test NFV deployment in distributed cloud-edge environment and implemented as a small-scale proof-of-concept system. As the name suggests, the system is aligned with NFV Management and orchestration (MANO) architecture¹⁴ and utilizes the edge resources in addition to the central data centers. Cziva et

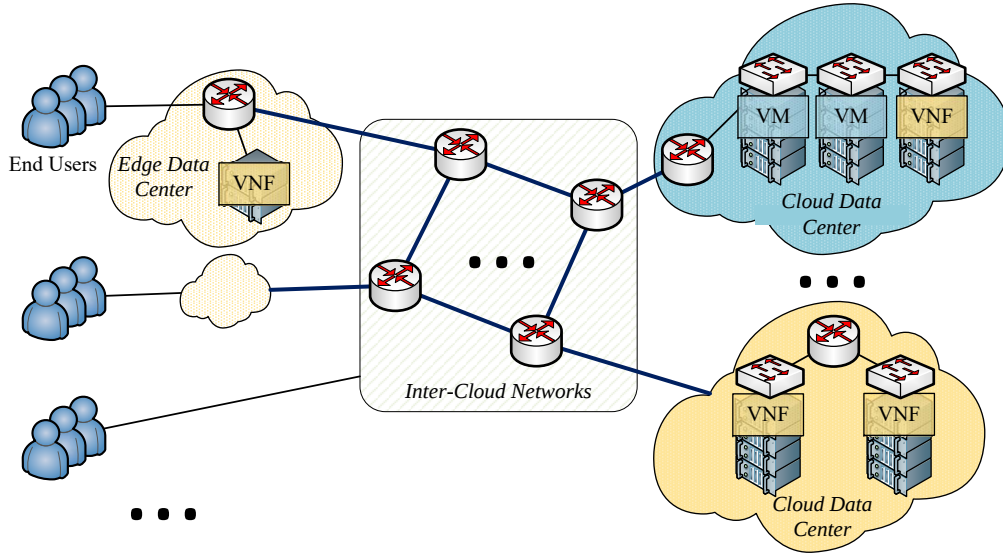


FIGURE 1 Architecture of NFV in edge-cloud environments.

al.⁷ proposed container-based network function framework supporting edge computing which can create a network function as a container and migrate to other sites in the edge of the network as like container migration. The proof-of-concept system is developed in Glasgow using OpenDayLight, OpenVSwitch, and Linux systems.

Although these practical systems provide the working example of the proposed approach, building such a system needs expensive equipment and time-consuming effort to setup the environment and implement the proposed policies. Also, it is difficult to evaluate the new approach in a large-scale with the proof-of-concept systems, as building the large-scale system is in need of even more efforts and resources. Therefore, an efficient and quick toolkit is still necessary for simulation and evaluation in a large-scale and cost-effective manner. The work proposed in this paper meets this requirement by developing software simulation toolkit for modeling and simulation of NFV and SFCs for their use in evaluating policies for edge and cloud computing environments.

3 | SIMULATION AND MODELING NFV IN EDGE AND CLOUD COMPUTING

In order to evaluate and test new algorithms regarding the NFV and SFC management as well as corresponding network functionalities, a scalable methodology has to be employed to see the effectiveness of the algorithm in a large-scale. It is impractical to test such experimental algorithms in a production environment due to unpredictable results from the new method. An experiment in a small-scale testbed is an alternative practical solution to validate the effectiveness, but the impact at large-scale can be hardly captured in the small-scale testbed system. Also, empirical implementation of the new algorithm to deploy onto the testbed can be time-consuming and challenging especially if it is in the early stage. Therefore, simulation has been widely adopted in science to simplify the evaluation with reasonable accuracy which can be compared with baselines under the same condition. With simulation, multiple experiments with various configuration and settings can be performed in a short time with automated scripts.

In the same way, simulation of NFV in edge-cloud environment is critical in order to reduce the evaluation time and to simplify the process. Ultimately, the accessible, adaptable, and scalable simulation toolkit will foster the innovation in NFV by reducing the effort of evaluating new methods and algorithms in the field. The innovative ideas can be implemented and evaluated with the simulation toolkit without having to spend enormous time preparing testbed infrastructure and deploying the new algorithm onto the system.

The overall architecture of the NFV implementation in edge-cloud environment is presented in Figure 1. Our simulation framework is based on the architectural components noted in the figure. Edge data centers are micro-scale data centers (e.g., a set of network routers, access points, or a base station with a few computing servers),

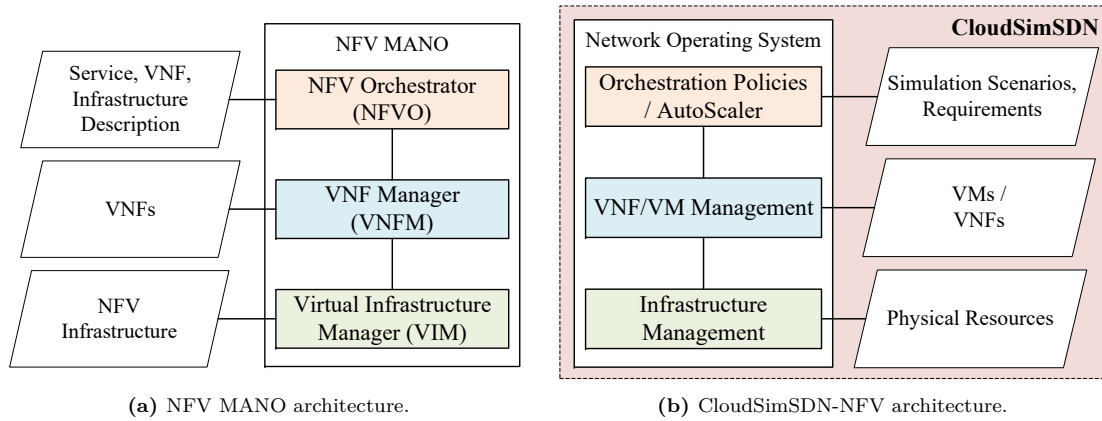


FIGURE 2 CloudSimSDN-NFV architecture aligned with NFV MANO.

located close to the end-users. End-users can access application services from edge data centers or cloud data centers through inter-cloud networks (i.e., backbone network). VNFs can be placed either in the edge or cloud data centers. Cloud data centers can provide VMs for application services and VNFs for network services. More detailed models and abstractions are described in the following subsections.

3.1 | Modeling NFV and Edge Computing

NFV Management and orchestration (MANO) architecture¹⁴ has been proposed by European Telecommunications Standards Institute (ETSI) and widely accepted in the communication field. MANO defines three main components for orchestration and management. NFV Orchestrator (NFVO) receives the business and operational requirements and coordinates VNFs in the system to provide the required services. NFVO can manage various VNFs through VNF Managers (VNFM) to actually create and provision VNFs inter-operating with SDN elements. Underneath VNFM resides Virtual Infrastructure Manager (VIM) which controls infrastructure resources hosting VNFs.

Our simulation framework aligns with MANO architecture to provide standardized NFV simulation and evaluation method as shown in Figure 2. We modeled the network infrastructure as physical resources, such as physical hosts, switches, and links. The modeled infrastructure is managed by the infrastructure manager (i.e., VIM) for network configuration, address assignment, and physical path. As VNFM can manage VNFs in the system, the simulation framework can create, delete, or migrate VNFs upon the infrastructure as VM instances. VNF management is decided by the orchestrator (i.e., NFVO) which is modeled as various policies in our framework. With the input requirements, scenarios, and virtual network configurations, user-specified policy in the simulation framework decides the orchestration policies among VNFs, e.g., increasing resources for a VNF, decreasing resources in case of under-utilization, or migrating to a different infrastructure.

In addition to NFV functionalities, we modeled and integrated edge computing simulation in the framework. Similar to the central cloud data center, an edge data center is modeled with a set of physical infrastructure (hosts, switches, and links) interconnected to each other. Edge data centers are differentiated with the limited resource capacity and the type of data center. For inter-cloud network between the edge and central clouds, inter-cloud switches and links are modeled to connect those data centers. Each data center can run its own network and computing resource management policies with separate inter-cloud network policies to connect among data centers. Location information on VMs, VNFs, and physical hosts is visible globally in order to simplify lookup and management for simulation purposes.

TABLE 1 Parameters for Service Function Chaining descriptors template

Type	Parameters						
Node	name	type	size	pes	mips	mipoper	ram
	datacenter	subdatacenters	host				
links	name	source	destination	bandwidth			
policies	name	source	destination	flowname	sfc	expected_time	

4 | DESIGN AND IMPLEMENTATION

The new simulation framework is designed and developed upon CloudSim¹⁵, a discrete event-driven simulation framework implemented in Java language. It follows the object-oriented programming model same as its fundamental CloudSim. Components in NFV and edge computing are designed as Java Classes, which can be extended or substituted based on the requirements and simulation scenarios according to the object-oriented model. In this section, we present the design and implementation details of each component in CloudSimSDN-NFV for simulating NFV and edge computing.

4.1 | Event-driven Simulation

In CloudSim, every simulation occurs by sending and receiving events between entities. For example, when a VM is created in a data center, a VM request event is sent to the data center entity. Then, the event is received by the data center which allocates resources for the requested VM using a VM allocation policy assigned to the data center. Similarly, a CPU workload can be sent to the VM through the data center entity. The processing scheduler of the VM will receive the workload submission event, calculate the workload end time, then send back the workload completion event with the calculated end time.

CloudSimSDN-NFV follows the same principle to simulate network transmissions, VNF creation and deletion, and inter-cloud events. It sends and receives events between entities, and the event delay is calculated by policies and schedulers which can be customized with various scheduling methods and algorithms.

4.2 | Virtualized Network Function (VNF)

VNF is designed and implemented by extending the VM Class from the original CloudSimSDN¹¹. Being virtualized network appliance with high CPU requirements, VNF is designed to have the same characteristic, such as processing capacity modeled as the number of cores (Processing Elements) and MIPS (Million Instructions Per Second), memory size, and storage size. In addition to those general VM specifications, VNFs have a specific field named MIPO (Million Instructions Per Operation), which models the throughput of the VNF. MIPO specifies the CPU workload length for a single network operation provided by the VNF, which can provide the throughput of the VNF along with MIPS. For example, a VNF with 1000 MIPS capacity and 10 MIPO can handle 100 requests (operations) per second throughput. MIPS is also used by VM/VNF allocation policies which decides a physical host to place the VM. If a host is lack of available MIPS for the requested VM/VNF, the allocation policy will try to allocate it in the other host.

4.3 | Service Function Chaining (SFC)

We offer the simulation capability of SFC in the framework (see Figure 3). SFC is defined as a chain of multiple service functions (i.e., VNFs) as an ordered and directed list. Enforcing SFC is determined in ServiceFunctionForwarder which is in charge of forwarding the matched packet to the chain of SFs. The forwarder checks every ServiceFunctionChainPolicy defined in the current simulation configuration, which is composed of the source and destination VMs. For example, if a network flow from VM1 to VM2 has to go through a chain of two VNFs, VNF1 and VNF2, then we can simulate such behavior by creating two VMs and two VNFs in the data center, defining a SFC with VNF1 and VNF2, and enforcing a policy to redirect the network flow from VM1 to VM2 to pass through the defined SFC. These SFCs and enforcing policies are defined in the simulation configuration, loaded and deployed when the

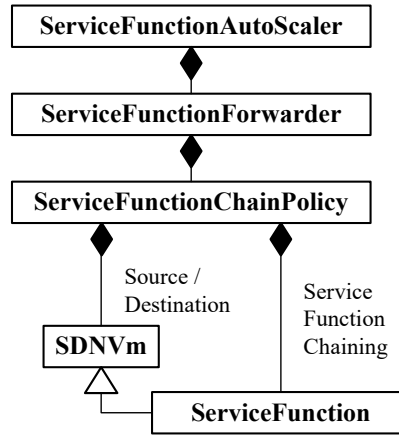


FIGURE 3 Class Diagram for simulating Service Function Chaining.

simulation started and enforced for all network traffics matching the policy. The simulation configuration of virtual topology (shown in Table 1) described in JSON language contains three categories: (1) Node template defines the name of the instance, type (VNF, VM, Container), resources requirement (disk size, CPU (pe), CPU speed (mips), Memory size (ram)), and packet processing speed MIPO for VNF Node type (mipoper), as well as optional allocation requirement for the instance (datacenter, subdatacenters, and host); (2) The template for virtual links defines the name of virtual links between source VM/container to destination VM/container with specific bandwidth as the option. If the name of the virtual link is "default", the packet scheduler will allocate residual bandwidth to the link after satisfying non-default links. The policy for packet processing can be changed accordingly based on the scenarios of the user simulation; (3) The third part of the template is the same as the VNF Forwarding Graph Descriptors (VNFFGD) used in the Tacker²² of OpenStack for creating forwarder and SFC. It defines the name of the SFC police, the source and destination of the SFC and which link should be enforced in the SFC (flowname), the directed VNF chaining sequence for the link (e.g., "sfc":["vnf1", "vnf2", "vnf4"]), as well as the expected QoS in terms of end-to-end delay for the link (expected_time). However, the users can also change and create new VMs/VNFs anytime during the simulation and change the connectivity as well as allocated resources according to the user-defined algorithms dynamically.

On top of the SFC policies and the forwarder, we implement the auto-scaling policy for SFCs. In every time interval, ServiceFunctionAutoScaler retrieves the average end-to-end latency for packets in the SFC and the utilization of SFs in the SFC. Based on the pre-defined auto-scaling policy, the auto-scaler can increase the capacity of the SF (vertical scale) and/or the number of SF machines for the specific function (horizontal scale). It can also increase the allocated bandwidth for the specific SFC if the bottleneck is at the network. We present a use case scenario of SFC auto-scaling in the evaluation section.

4.4 | Packet Scheduler

We model the packet scheduler similar to the Cloudlet scheduling in the CloudSim, as shown in Figure 4. In the original CloudSim, a computing workload for CPU processing is modeled as a Cloudlet which has the length of the processing workload. CloudletScheduler is in charge of scheduling the processing workload in each VM based on the simulation scenario. In CloudletSchedulerTimeShared, the processor capacity is evenly shared by all Cloudlets submitted and currently processed at the VM. For example, if five Cloudlets are submitted to a VM, the CPU capacity of the VM is shared among them so that each Cloudlet can be assigned 20% of the total CPU capacity. On the other hand, CloudletSchedulerSpaceShared processes only one Cloudlet at each time so that 100% of the CPU capacity will be assigned to the first Cloudlet submitted to the VM. The other Cloudlets are in the waiting list and processed in the queue once the earlier Cloudlet completes the processing.

In CloudSimSDN-NFV, we design the network packet scheduler with Packet and PacketScheduler Classes similar to Cloudlet and CloudletScheduler. Packet Class represents the network transmission workload which has the size

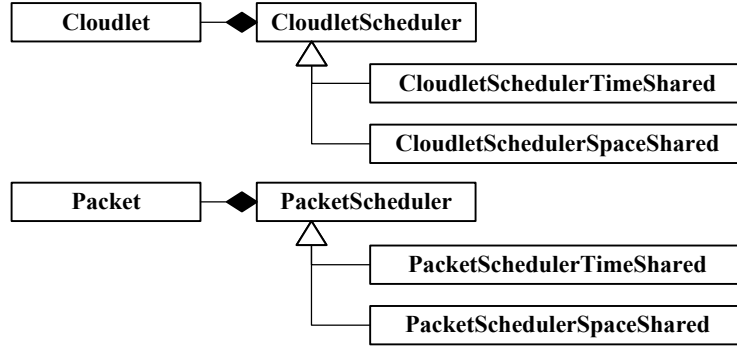


FIGURE 4 PacketScheduler to simulate network packets along with CloudletScheduler of the original CloudSim.

of the network packet. PacketScheduler distributes the available network bandwidth among currently transferring Packets with the same source and destination VMs. If multiple flows share the same physical link, the bandwidth of the physical link is distributed among these flows¹¹ and then PacketScheduler can allocate the distributed bandwidth onto Packets in the scheduler. Similar to CloudletScheduler, we implement PacketScheduler with two models, time-sharing and space-sharing. In time-sharing, the available bandwidth is equally shared among Packets from the same source VM to the same destination VM. In SpaceShared, the entire bandwidth of the virtual network is allocated to the first Packet submitted to the network, and the rest are waiting in the queue until the transmission of the first Packet is completed.

4.5 | Edge Computing

For simulating edge computing, we regard the edge resource as part of an edge data center which has the similar characteristic of a cloud data center. For example, the edge data center will have physical hosts with limited capacity and switches to connect the hosts and end-users. The difference between edge and cloud data centers will be the size and scale of the resources provided in the data center. In terms of simulation, an edge data center and a cloud data center can be regarded as the same entity, while the edge data center would provision less amount of computing resources compared to the cloud data center. Other than the resource capacity, both edge and cloud data centers can have its own computing resource management policies and network rules. Thus, we use the data center Class implemented in the original CloudSimSDN to represent both edge and cloud data centers.

However, simulating edge and cloud data centers in the same scenario has to consider running multiple data centers simultaneously. The original CloudSimSDN was potentially capable of supporting multiple data centers in its design, but it was not implemented properly. In this work, we consider the multi-cloud data center scenario case and implement the correlated entities to work without interfering objects of other data centers. The new framework can create multiple data centers with different configurations, e.g., an edge data center with a few low-capacity hosts, and a large-scale data center with thousands of high-performance hosts. The network configurations can be independently set up to reflect the real-world scenario.

Upon properly implementing multiple data centers, inter-cloud networks have been designed and implemented in this work to support network traffics between data centers. We extended the original Switch Class to InterCloudSwitch Class in order to represent backbone switches connecting distributed data centers. Additionally, GatewaySwitch Class is introduced to indicate gateway switches in a data center. Network traffics from one data center to another have to pass through the gateway of the origin data center, the distributed inter-cloud switches, the gateway of the destination data center, and then to the destination host. Various network topologies can be defined individually for different data centers and inter-cloud networks. For example, an edge data center can be defined with a canonical tree topology, a cloud data center with fat-tree, and the inter-cloud network with mesh topology. Furthermore, as the SDN routing elements are separated from the physical elements in the principle of SDN, every hosts and switches are considered as a Node Class extension for virtual routing. Therefore, data centers can also be defined with the hybrid or server-centric topology, such as, BCube, Dcell, and hypercube, without switches.

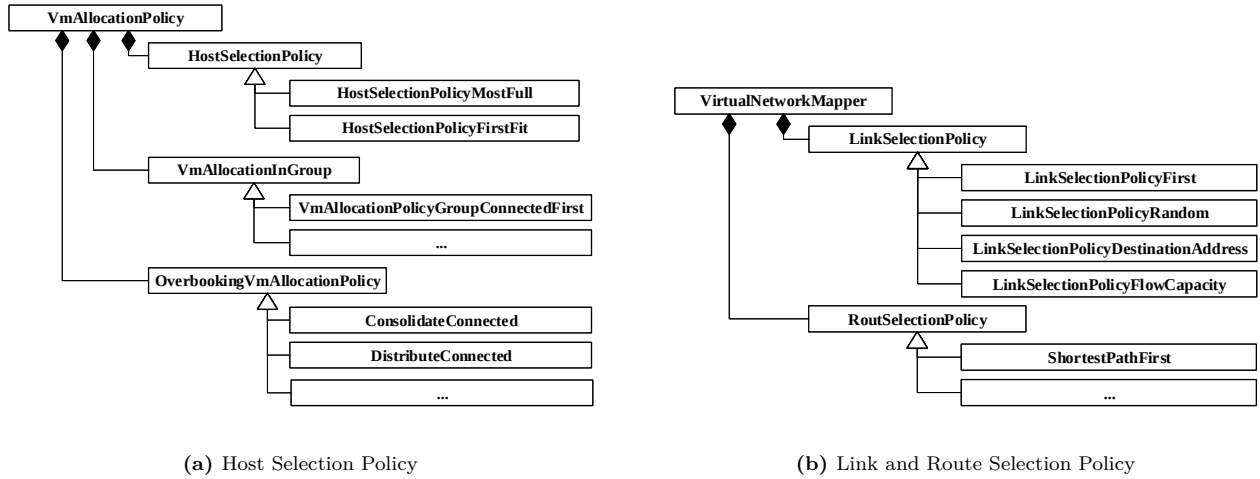


FIGURE 5 Allocation and network traffic engineering policies implemented in the framework.

TABLE 2 Basic parameters supported for simulation

Type	Parameters						
computing	CPU	Memory	Disk	Workloads	Task Scheduling	Task Priority	Overbooking Ratio
networking	Bandwidth	Topology	Switch Buffer	Ports	Channel Priority	Control Plane	Data Plane
monitoring	Statistic	Energy Consumption	Utilization	Response time	Network Delay	Fault Handling	
virtualization	Connectivity	Allocation	Lifecycle	Flow Scheduling	QoS	SLA	

4.6 | Customizable Policies

In CloudSimSDN-NFV, we implemented the functionalities to evaluate and simulate the SDN-NFV-enabled edge and cloud computing in terms of computing, networking, monitoring, and corresponding virtualization (shown in Table 2). Policies, virtualization technology, and algorithms can be implemented and evaluated based on the optimization of these parameters, such as dynamic flow rerouting based on the processing speed of VNF and end-to-end delay along the SFC, and optimizing the allocation for the new replicated VNFs during auto-scaling based on the connectivity, network topology, and available bandwidth, etc. Based on the current simulation framework, we have evaluated the proposed latency-aware VNF dynamic provisioning algorithm to utilize both edge and cloud resources²³ and the unified framework of different auto-scaling algorithms to minimize the cost with the guarantees of SLA and end-to-end delay along the SFCs²⁴.

Beside the auto-scaling policies discussed in Section 4.3, we modularize the VM/VNF/Container allocation and network mapping policies for simple customization as shown in Figure 5. Note that we don't show all algorithms that we have implemented but the basic template policies. VmAllocationPolicy defined in the original CloudSim is extended with the separated HostSelectionPolicy which can select the best host to allocate a VM among hosts in the candidate list which sorting by the optimization policy. We implement two policies as the basic template for selecting a host: first fit and most full methods. First fit is to return the first host appeared in the list which has enough resources available for the requested VM. The most full method is to find the most occupied host in the list which can still provide the requested resources. As shown in Table 2, according to the objectives of optimization, the sequence of host list can be based on other parameters, such as computing, networking, and monitoring. Based on the connectivity among VNFs, VMs, and containers, we also implement policy VmAllocationPolicyGroupConnectedFirst to showcase the potential extensions. Furthermore, we implement more complex allocation policies OverbookingVmAllocationPolicy for the beginning and consequent dynamic resource management based on the CPU, memory, disk, and bandwidth overbooking by utilizing migration. Therefore, other policies can be easily implemented by following the existing source code with minimal extension effort.

For a virtual network, VirtualNetworkMapper Class is in charge of mapping the requested virtual network topology onto the physical network elements. Similar to HostSelectionPolicy, a separate LinkSelectionPolicy Class implements which link to choose among multiple links in the physical topology. We implement several methods in the framework as an example. LinkSelectionPolicyFirst selects the first link in the candidate list regardless of the network capacity. Similarly, a random link can be selected with the Random policy. More practically, LinkSelectionPolicyDestinationAddress looks up the destination address of the network and assigns a certain link calculated by a modulo operation. For example, if there are two available links, the destination address with the odd number selects the first link, and the even number selects the second link. LinkSelectionPolicyFlowCapacity selects the link with the most remaining capacity, such as available bandwidth or remaining buffer in switches. It checks currently occupied network capacity of the links in the candidate list and returns the most empty one among them. Thus, the network transmission can be evenly distributed among multiple paths. LinkSelectionPolicy is also readily customizable to implement a new method. Moreover, as shown in Figure 5b, we also implement the policies for proactive flow routing in SDN virtual network. Based on both global physical and virtual network topology provided in the CloudSimSDN-NFV, we can proactively push every specific virtual or physical routing into the forwarding tables of every Node, such as based on the shortest path first algorithm using the global available bandwidth. Furthermore, researchers can also extend the proactive traffic engineering algorithm, e.g., the simulation of the policies for network security and jurisdiction: some flows must through to specific groups of servers within the jurisdiction and cross-region data flows. In general, for researchers who want to evaluate new approaches in VM allocation and/or network mapping, they can modify these policies with the proposed method and compare with the baselines to evaluate the effectiveness of the new approach.

4.7 | Energy Consumption and Utilization Monitor

The new framework also includes monitoring components for energy consumption and resource utilization. A Java interface is defined to record the utilization history of each element (e.g., a switch, a host, and a VM). After the pre-configured monitoring interval time, the monitoring event is triggered to measure the amount of processed workload and calculate the utilization of the element for the previous time period. At the end of the every monitoring event, the same event with the monitoring time interval is sent again so that the utilization history is recorded periodically. In addition to recording utilization, the estimated energy consumption for the time interval is also calculated with a linear model. In the current version, we used the linear model to estimate power consumption for hosts²⁵ and for switches²⁶ which can be replaced with other models upon the simulation scenario and requirement.

5 | EVALUATION USING TWO USE CASES

We evaluated the proposed simulation framework with two use case scenarios for NFV and edge computing. This section discusses and compares the effectiveness of different policies in NFV with two use case scenarios. In the first scenario, we evaluate the end-to-end delay and the estimated cost with different auto-scaling policies of SFC. We implement the auto-scaling policy for processing capacity of VNFs and bandwidth capacity for the chaining upon the simulation framework to compare the measured delay and resource usage. In the second scenario, the edge data center with limited resources is created to host some VNFs. We implement two VNF placement policies to choose between edge and cloud data centers: utilizing edge nodes for placing VNFs, or all VNFs placing in the cloud. These two use case scenarios present the potential usage of the simulation framework in various NFV and edge computing research. For both scenarios, we generated synthetic workloads based on Wikipedia trace and following the three-tier application model²⁷.

5.1 | Experiment 1: SFC Auto-Scaling Policies in NFV

The first use case is to evaluate different SFC auto-scaling policies within a single cloud data center. In this scenario, we created a large-scale cloud data center with 128 physical machines connected through 32 edge switches, 32 aggregation switches, and 16 core switches, forming 8-pod fat-tree²⁸ network topology. Thus, each pod consists of 4 edge switches, 4 aggregation switches, and 16 physical machines where 4 machines are connected to each edge switch.

TABLE 3 Configurations of different instances in experiments

Policy	Type (VNF/VM)	Mem (GB)	CPU (cores)	Disk (GB)	Bandwidth (Bytes/s)	Number of VNFs/VMs
NoScale-Max	lb/ids/fw	8	10/12/16	8	2,000,000	1/3/3
NoScale-Min	lb/ids/fw	8	2/6/8	8	500,000	1/1/1
Flavor	web/app/db	256	8/4/12	1000	1,500,000	8/24/2

Each physical machine is configured with 16 cores, and each core has 10,000 MIPS capacity. The network bandwidth of all links between switches and physical machines are equally set to 200 MBytes/sec.

5.1.1 | Virtual Topology and SFC Configuration

For virtual topology, we simulate three-tier web applications consisting of web, application, and database servers to form an entire application. There are 8 web servers, 24 application servers, and 2 database servers. The configuration of different types of servers are shown in Table 3. All servers are created as VMs in the simulation which receives workloads consisting of CPU processing and network transmission. In order to evaluate NFV functionalities, we prepare 4 VNFs from three different types of VNF Firewall (fw): VNF1, Load Balancer (lb): VNF2, VNF3, and Intrusion Detection System (ids): VNF4 as shown in Table 3, which are enforced by 4 different types of SFC policies. Furthermore, we assign MIPO to Load Balancer, IDS, Firewall as 20, 200, and 800 respectively. Network traffics between the VM servers are forwarded to different SFCs consisting of a subset of the 4 VNFs as shown in Table 4. As all network traffic between these VMs is diverted to transmit through the SFCs, these four VNFs can be a bottleneck of the application performance if the initial capacity of the VNFs is not enough to serve the entire application.

5.1.2 | Auto-scaling Policies

In this scenario, we evaluate three algorithm (*NoScale-Min*, *NoScale-Max*, and *AutoScale*) for different SFC auto-scaling policies in combination with two VM/VNF allocation policies (*LFF*, *MFF*). The configuration of *NoScale-Min* and *NoScale-Max* of different types of VNFs are shown in Table 3. In *NoScale-Min*, the auto-scaling feature has turned off and the minimal resources are allocated to the VNFs for the entire experiment duration. As VNFs have allocated the least amount of resources without any scaling, they are always in lack of resources which makes the network packets delayed to wait in the queue at VNFs. On the other hand, by using more CPU, bandwidth, and the same type of VNF in one SFC police, *NoScale-Max* policy allocates more than enough resources to VNFs from the beginning of the experiment without auto-scaling. As a large amount of resources are allocated for VNFs from the beginning, they can be underutilized if the submitted workload is less than estimated. In *AutoScale* policy, the initial resource allocation is the same as *NoScale-Min*, which allocates the minimal resources at the beginning of the experiment. However, it continuously monitors the utilization of VNFs and increases the size of the VNF when the utilization level is over the defined threshold. In this experiment, we set the threshold at 70% for auto-scaling. If it is possible to increase the capacity to twice of the current size for scaling-up in the current host (double the capacity

TABLE 4 SFC policies enforcing network transmission to divert to VNFs.

Source	Destination	Subset of VNFs
Web server	Application server	{VNF1, VNF2}
Application server	Database	{VNF3, VNF4}
Database	Application server	{VNF4, VNF3}
Application server	Web server	{VNF2}

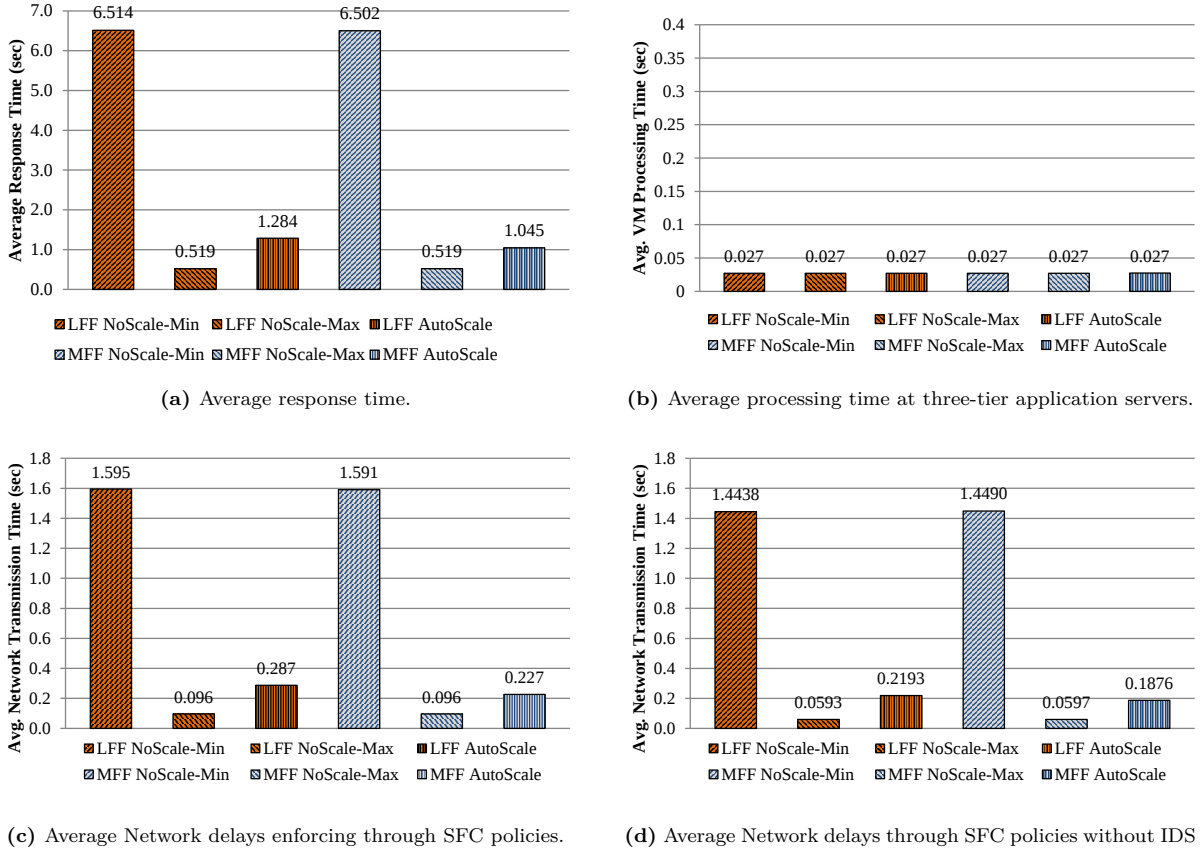


FIGURE 6 Average response and processing time with various auto-scaling policies.

of CPU and allocated bandwidth in this case), then scaling-out policy by duplicating the VNF of the same type is not an option. If not, we create another VNF of the current one and allocate it according to the LFF or MFF policy.

In addition to different auto-scaling policies, we evaluate the effectiveness of different VM and VNF allocation policies. *LFF* (Least Full First) algorithm finds the least full physical machines with less amount of allocated resources for a VM. In this algorithm, VMs and VNFs are dispersed across the data center, as a physical machine hosting a VM has less priority than a machine without any hosted VM. In contrast, *MFF* (Most Full First) algorithm selects the most allocated machine which has enough capacity to host the VM. The selected physical machine still has enough resources for the VM, but more resources are already allocated to the other VMs compared to the other candidate machines. VMs are consolidated into a smaller number of physical machines with this algorithm, which can reduce the number of machines and the energy consumption of the data center.

Furthermore, in scenarios where IDS is optional, we evaluate the influence of SFC length and the VNF processing delay on the end-to-end response time. By removing VNF4 from the SFC policies shown in the Table 4, the total length of SFC policies is changed from 7 to 5. We combine three auto-scaling policies with two VM/VNF allocation policies to create six cases and present in the following.

5.1.3 | Evaluation Results

At first, we measure the response time of all workloads submitted to the application as shown in Figure 6. The average end-to-end response time is depicted in Figure 6a. As expected, *NoScale-Min* policies in both LFF and MFF result in exceptionally longer response time due to the lack of resources in VNFs. On the other hand, the delay is short in *NoScale-Max* policy as the sufficient resources are allocated to VNFs. In *AutoScale* policy, the response time is slightly longer than *NoScale-Max* but not as long as *NoScale-Min*, since the highly utilized VNFs can acquire more resources by the auto-scaling policy after the short period of time. The observed difference is not significant between

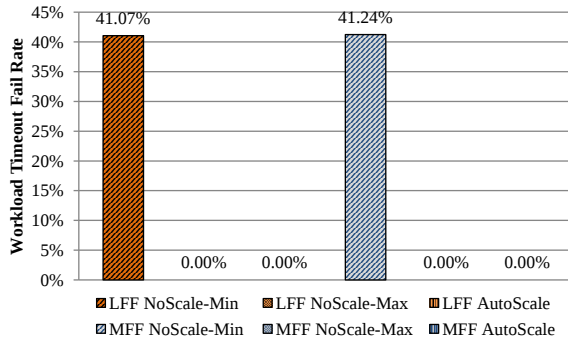


FIGURE 7 Workload failure timeout rate.

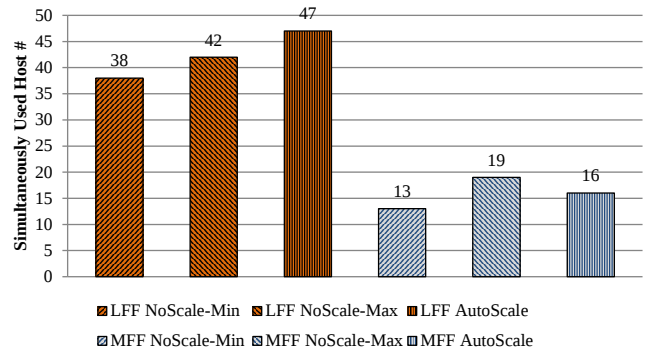


FIGURE 8 Simultaneously used hosts number.

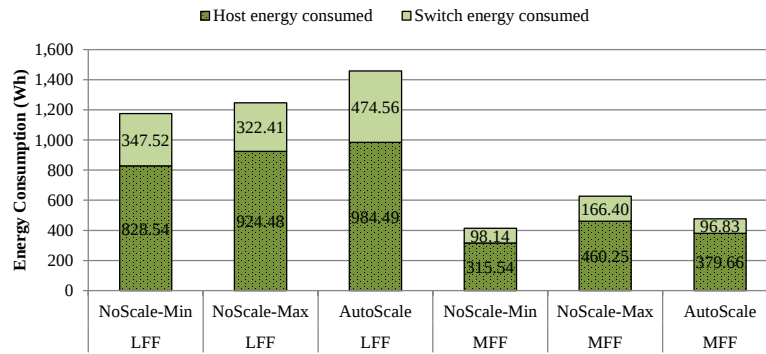


FIGURE 9 Energy consumption in hosts and switches in the data center.

LFF and MFF in every auto-scaling policy. To further explain the average response time (end-to-end delay), we separate it into processing time of application servers and network delays along SFCs.

Figure 6b and 6c respectively show the average VM processing time at the service application servers (web, application, and database VMs) and the average network transmission time including delays at VNFs enforced in SFC. We can observe the the application processing time at VM servers remains same across all six combinations which explain the processing power for web, application, and database servers are enough to process all workloads within short time. However, the network transmission time has a significant difference between different policy combinations due to the delay in VNFs occurred by SFC enforcement. As the bandwidth is much sufficient for the services, it is the VNF delays along the SFCs contributes the overall end-to-end delay the which includes packet processing time and buffer waiting time. In *NoScale-Min*, the initial resources allocated for VNFs are insufficient to process the network traffic which results in the packets delayed in SFC. *NoScale-Max* and *AutoScale* policies reduce the network transmission time significantly as the resources for VNFs are sufficient from the beginning (*NoScale-Max*) or increased to be sufficient after a while (*AutoScale*). In the case of *AutoScale*, the average network delay of LFF allocation policy is larger than the MFF. The LFF policy allocates new replicated VNF to new hosts. As a result, a more distributed allocation leads to a larger network transmission time due to more network hops. Furthermore, in the scenario of no IDS in SFCs, Figure 6d indicates that the network delays along SFCs reduced accordingly as the result of shorter SFC length.

We also measured the network timeout rate in each policy, which leads to the workload failure. The timeout threshold is set to 5 seconds so that a workload taking longer than 5 seconds is dropped and counted as timeout packet. Figure 7 depicts the timeout rate among all workloads with different policies. *NoScale-Min* policy results in dropping over 40% workload by timeout, whereas none of the workloads is dropped in the other two policies.

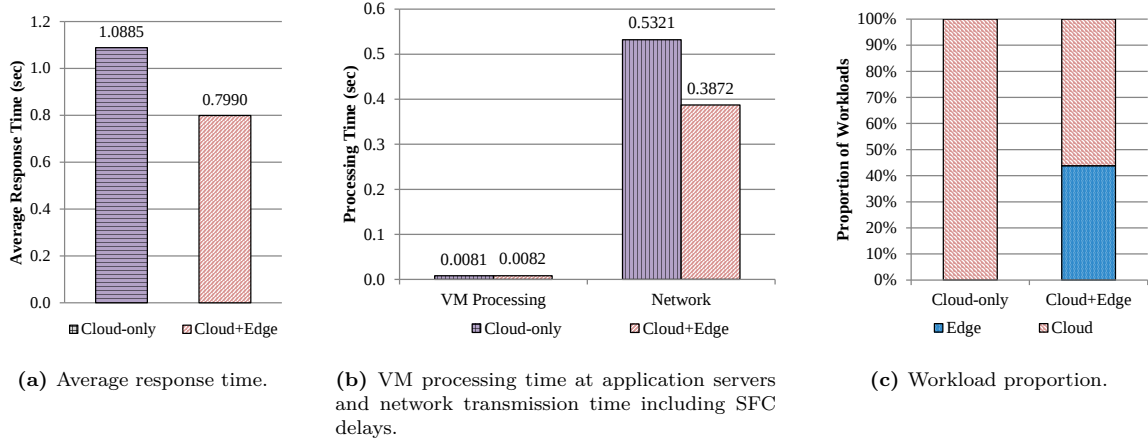


FIGURE 10 Performance evaluation results of VNF placement policies in edge-cloud environment.

Then we measured the maximum number of simultaneously used hosts. As shown in Figure 8, overall the LFF algorithm uses more hosts. For *AutoScale* SFC algorithm, *MFF* is better policy compared to the LFF which fully utilizes the computing resources. Although the simultaneously used hosts' number of *MFF AutoScale* is bigger than the *MFF NoScale-Min* as necessary to alleviate the workload failure, it is smaller compared to the *MFF NoScale-Max*. By using fewer hosts, the *MFF AutoScale* maximizes revenue generation for service provider and also offers cost saving for cloud users.

Finally, we measure the estimated energy consumption in the data center for hosts and switches, as shown in Figure 9. A clear difference can be seen between LFF and MFF, as LFF consumes more energy in any auto-scaling policies compared to MFF. For example, LFF with *AutoScale* consumed 984.49 Wh in hosts and 474.56 Wh in switches totaling 1,459.05 Wh, whereas MFF with *AutoScale* consumed 476.49 in total (379.66 in hosts plus 96.83 in switches). In MFF, more numbers of VMs are allocated in a single host which leads to consolidating VMs and workloads into less number of hosts. The unused hosts can be turned to the idle mode which can save power consumption significantly. Interestingly, more power was consumed in *AutoScale* compared to *NoScale-Max* in LFF allocation policy, because more numbers of smaller VNFs were created in *AutoScale* due to LFF policy. In LFF, VNFs are distributed across the data center leading to utilizing more numbers of physical hosts. The duplicated VNFs created from the auto-scaling policy after the initial placement are also dispersed leading to utilizing more physical hosts to be used. Also, these scattered VNFs increase the volume of network traffic in the system, which results in more energy consumption in switches.

5.2 | Experiment 2: NFV Placement Policy in Edge-Cloud Environment

In the second experiment, we evaluate different VNF placement policies in edge-cloud environments. The composition of various service components is a crucial technique to optimize the resource allocation and utilization of edge and cloud resources, as well as to reduce end-to-end latency for end-users. A cloud data center is built with an excessive amount of resources to provide infrastructure to serve vast applications and customers, whereas edge data centers are equipped with limited resources such as a base station or a wireless access point. Therefore, the limited number of VNFs can be placed in edge data centers although the edge resources can help to reduce end-to-end delays for end-users as the network traffics can be processed closer to the users.

In this experiment, we created an edge and a cloud data center (i.e., network data center) to place VNFs, in addition to the extra data center (i.e., application data center) to host application servers. The edge data center has limited capacity with 4 hosts (each with 16 cores and 10,000 MIPS/core), whereas two cloud data centers (one for application servers and another for network functions) have the very large amount of resources to host an almost infinite number of VMs and VNFs. End-users are placed within the edge data center. The network latency of a link within the same data center is set to 1 msec. Edge and cloud data centers are connected through four inter-cloud

switches linearly connected to each other with the network latency set to 100 msec. Thus, a network latency for inter-cloud traffic can be up to 400 msec, while less than 10 msec within the same data center. In this experiment, we use a canonical tree topology to simplify the creation of various data centers.

The application is composed of 12 servers running as VMs and placed in the application data center, which is accessed by 4 end-users created in the edge data center. All network traffics from the end-users to the application servers are enforced to go through VNF1, and the traffics of opposite direction are through VNF2. In this experiment, the auto-scaling policy is always turned on, so that the number of VNFs to serve the same type of function can be increased if the utilization of the VNF is over the threshold, same as the auto-scaling policy experimented in the previous scenario. When a VNF is duplicated to serve more network traffics, the new VNF can be placed either in edge or clouds. We use different policies to select where to place the duplicated VNFs.

5.2.1 | VNF Placement Policies

We implemented two different placement policies in the simulation to select a place for VNFs. *Cloud-only* policy is to utilize only network cloud data center for placing VNFs and no edge resources. As cloud data centers can provide large capacity, this policy does not consider the resource requirement of the VNF and places all VNFs in the cloud. On the other hand, *Cloud+Edge* policy considers the edge data center as the option, so that if there are available resources in the edge, the new VNF is placed in the edge first. If the edge data center has insufficient resources, the cloud data center will host the VNF. Note that the edge data center itself has limited resources insufficient to place the required VNFs for the whole workload, therefore the edge-only policy was not considered in the experiment.

5.2.2 | Evaluation Results

Figure 10 shows the average response time, CPU and network processing time, and the proportion of workloads processed in each data center, measured in the simulation. Average application response time is measured from the time that end-users send a request to the application server via VNFs until the response arrives at the end-users. Both request and response are enforced to forward through VNFs. As depicted in Figure 10a, the average response time is reduced by 26.6% in Cloud+Edge policy from 1.0885 seconds with Cloud-only policy to 0.7990 seconds, thanks to the reduced network latency by placing VNFs in edge resources closer to end-users. VM processing time in application servers is measured similarly at approximate 0.008 seconds in both policies, but the network transmission time has a significant difference between two policies (see Figure 10b). This result shows that the reduction of network transmission time in Cloud+Edge policy is the reason for the improved response time. The last figure (10c) shows the proportion of VNF workloads processed in edge and cloud data centers. In Cloud-only policy, all packets are obviously processed only in the network cloud data center, whereas about 45% packets are processed in the edge data center with Cloud+Edge policy. The rests (55% of workloads) are still processed in the network cloud, due to the lack of resources in edge data centers for the newly duplicated VNFs by auto-scaling policy.

6 | FUTURE DIRECTIONS AND EXTENSIONS

As an event-driven simulation framework, CloudSimSDN-NFV supports evaluation of both networking and computing resource management in the context of edge and cloud computing. For the networking processing, we have already evaluated the accuracy compared with Mininet¹¹. However, when it comes to the scenarios that require high precision evaluation and prediction, such as low level infrastructure-related parameters, it depends on the profiling and modeling. By profiling the real data and identifying the acquired parameters, we can translate the mathematical models to the simulation framework with higher precision. In this section, we discuss our thoughts on how CloudSimSDN-NFV can be extended comprehensively to support: (1) OpenFlow-like Capabilities, (2) live migration evaluation, and (3) energy modeling.

6.1 | OpenFlow-like Capabilities

Current CloudSimSDN-NFV enables several SDN functionalities such as dynamic proactive packet routing, forwarding a packet to set of ports, and modifying the destination of a packet or dropping the packet. As OpenFlow²⁹ is the de facto SDN protocol, it can facilitate the functionalities through proactive and reactive flows based on packet-in and packet-out between controller and switch. For OpenFlow packet-in and packet-out simulation, we can extend the current class of packet and specify the packet label for the OpenFlow protocol. Furthermore, based on the event-driven simulation principles, several events according to the user-defined algorithms and policies can be created in the class of *SDNHost*, *SDNDataCenter* and *NetworkingOperationSystem*. These features can be used in evaluating new algorithms for placement of distributed SDN controllers. Such algorithms aim at minimizing the communication delays of packet-in/packet-out messages between SDN-enabled switches and the SDN controllers.

6.2 | Comprehensive Modeling of Live Migration

Generally, existing works use live migration as an approach for dynamically allocating VMs^{12,8} to realize the objectives such as optimal end-to-end delay and energy consumption. However, live migration can significantly influence the QoS and SLA by using both computing and networking resources. The current framework supports the mechanism of live migration which allows instance (VNF/VM/container) stops on the source host and resumes at the destination host with available resource checking. Nevertheless, it does not consider the emulation model of iterative dirty memory copy through network transmission during the migration. It only pauses the processing of workload in the *CloudletScheduler*, and reroutes the packet to the new destination when the migration is finished. For a comprehensive live migration model, we can extend it based on the different phases and the prediction modeling³⁰. That is, we can define the migration behaviors in the event processing functions of the *SDNDataCenter* and *NetworkingOperationSystem* according to the live migration phases: pre-migration, dirty page synchronization, stop-and-copy, and post-migration phases³⁰. For example, with the support of the current Channel manager and Cloudlet scheduler, it pauses the packet transmission and workload processing within the migrating instance during the downtime of live migration. Therefore, the overheads and performance of live migration such as live migration time, downtime, and the amount of transferred data can be evaluated. On top of the live migration modeling, we can schedule the migration flows and multiple live migration planning with the global network topology.

6.3 | Energy Modeling

Researchers who focus on the performance comparison between the current energy models and the state-of-art model of NFV can use CloudSimSDN-NFV directly to produce relatively accurate results. Because NFV as the software is still using the same networking resources (switches)²⁶ and computing resources (physical hosts in cloud data centers)²⁵. Power modeling of the virtualized data center can be categorized into analytic power estimation modeling²⁵ and empirical measurement-based characterization³¹. For comprehensive energy modeling, many issues such as (1) power consumption model for other devices like storage and memory; and (2) temperature flow model between different components and servers. Based on the currently available monitoring ability and extension flexibility, many researchers have proposed new solutions for energy modeling. Louis et al. extend the CloudSim to simulate the storage power consumption in the data centers³². Ilager et al. also extended the components of the server energy model in CloudSim by taking runtime temperature into consideration³³. They modeled the hot spots generated by the running physical hosts which have a significant impact on the energy cost for the whole cooling and power solution of the cloud and edge computing. Furthermore, researchers can also implement other power consumption models for VM/VNFs/container as mathematical models³⁴ in virtualized environments and OpenFlow switches³⁵ for SDN network. In summary, future directions for energy modeling of NFV and other virtualized features in the cloud and edge computing should integrate the energy models of memory, disk, and cooling by characterizing the low-level software behavior.

7 | SUMMARY AND CONCLUSIONS

The advancement of virtualization technology and exploding computing capacity lead to several innovative paradigm shifts in computing and networking industry, including the emergence of cloud computing, NFV, SDN, and edge computing. Cloud computing and SDN have been studied in the research community for a decade and widely accepted in the industry, while NFV and edge computing are still in the early stage of research and implementation in real-world. Although many state-of-the-art works are presented in the literature, the simple and quick evaluation framework can foster the advancement exponentially.

In this paper, we propose CloudSimSDN-NFV, a new simulation framework to evaluate NFV functionalities in edge and cloud computing along with other SDN functionalities and cloud computing environments. The framework is designed and developed upon CloudSimSDN, which is built on the well-studied CloudSim toolkit. We described the modeling and simulation of NFV and edge computing and the detailed design and implementation of our framework. Two use case scenarios are presented to help understand how to use the new tool, and several algorithms are implemented and evaluated upon the framework. The results show that our framework can be efficiently exploited for quick evaluation of various approaches in the simulation.

To improve the proposed framework, more in-built policies can be added and implemented for VNF placement and SFC composition. For example, a network topology-aware VNF policy can be implemented and tested within the simulation, as well as an SFC composition policy working with multiple network data centers. Under the proposed simulation framework, we elaborated the potential extensions and future directions in respect of OpenFlow-like capabilities, comprehensive live migration modeling, and energy modeling. Therefore, we expect that our simulation framework can empower researchers in conducting advanced investigations in NFV, edge, and cloud computing to foster new innovations.

SOFTWARE AVAILABILITY

This software is released in open source as CloudSimSDN v2.0. The code can be downloaded from <https://github.com/Cloudslab/CloudSimSDN>.

ACKNOWLEDGMENTS

This work is partially supported by an ARC Discovery Project. We thank Huaming Wu, and Adel Nadjaran Toosi for the valuable discussion and feedback to improve the quality of the simulation framework.

References

1. Buyya R, Yeo C. S, Venugopal S, Broberg J, Brandic I. Cloud computing and emerging IT platforms: Vision, hype, and reality for delivering computing as the 5th utility. *Future Generation Computer Systems*. 2009;25(6):599–616.
2. Son J, Buyya R. A Taxonomy of Software-Defined Networking (SDN)-Enabled Cloud Computing. *ACM Computing Surveys*. 2018;51(3):59:1–59:36.
3. European Telecommunications Standards Institute (ETSI) . *Network Functions Virtualisation*. 2018.
4. Chang C, Srirama S. N, Buyya R. Internet of Things (IoT) and New Computing Paradigms. In: *Fog and Edge Computing: Principles and Paradigms*. New York, USA: Wiley Press 2019.
5. Boubendir A, Bertin E, Simoni N. On-demand, dynamic and at-the-edge VNF deployment model application to Web Real-Time Communications. In: *Proceedings of 12th International Conference on Network and Service Management (CNSM)*:318-323; 2016.

6. Dominicini C. K, Vassoler G. L, Meneses L. F, Villaca R. S, Ribeiro M. R. N, Martinello M. VirtPhy: Fully Programmable NFV Orchestration Architecture for Edge Data Centers. *IEEE Transactions on Network and Service Management*. 2017;14(4):817-830.
7. Cziva R, Pezaros D. P. Container Network Functions: Bringing NFV to the Network Edge. *IEEE Communications Magazine*. 2017;55(6):24-31.
8. Cziva R, Anagnostopoulos C, Pezaros D. P. Dynamic, Latency-Optimal vNF Placement at the Network Edge. In: Proceedings of IEEE INFOCOM:693–701; 2018.
9. Mijumbi R, Serrat J, Gorricho J, Latre S, Charalambides M, Lopez D. Management and orchestration challenges in network functions virtualization. *IEEE Communications Magazine*. 2016;54(1):98-105.
10. Lingen F, Yannuzzi M, Jain A, et al. The Unavoidable Convergence of NFV, 5G, and Fog: A Model-Driven Approach to Bridge Cloud and Edge. *IEEE Communications Magazine*. 2017;55(8):28-35.
11. Son J, Dastjerdi A. V, Calheiros R. N, Ji X, Yoon Y, Buyya R. CloudSimSDN: Modeling and Simulation of Software-Defined Cloud Data Centers. In: Proceedings of the 15th IEEE/ACM International Symposium on Cluster, Cloud and Grid Computing:475-484; 2015.
12. Son J, Dastjerdi A. V, Calheiros R. N, Buyya R. SLA-Aware and Energy-Efficient Dynamic Overbooking in SDN-Based Cloud Data Centers. *IEEE Transactions on Sustainable Computing*. 2017;2(2):76-89.
13. Son J, Buyya R. Priority-aware VM Allocation and Network Bandwidth Provisioning in Software-Defined Networking (SDN)-enabled Clouds. *IEEE Transactions on Sustainable Computing*. 2019;4(1):17-28.
14. European Telecommunications Standards Institute (ETSI) . *Open Source NFV Management and Orchestration (MANO)*. 2018.
15. Calheiros R. N, Ranjan R, Beloglazov A, De Rose C. A, Buyya R. CloudSim: a toolkit for modeling and simulation of cloud computing environments and evaluation of resource provisioning algorithms. *Software: Practice and Experience*. 2011;41(1):23–50.
16. Garg S. K, Buyya R. NetworkCloudSim: Modelling Parallel Applications in Cloud Simulations. In: Proceedings of the Fourth IEEE International Conference on Utility and Cloud Computing:105–113; 2011; Washington, DC, USA.
17. Piraghaj S. F, Dastjerdi A. V, Calheiros R. N, Buyya R. ContainerCloudSim: An environment for modeling and simulation of containers in cloud data centers. *Software: Practice and Experience*. 2017;47(4):505-521.
18. Gupta H, Dastjerdi V, Ghosh S. K, Buyya R. iFogSim: A toolkit for modeling and simulation of resource management techniques in the Internet of Things, Edge and Fog computing environments. *Software: Practice and Experience*. 2017;47(9):1275-1296.
19. Filho M. C. S, Oliveira R. L, Monteiro C. C, Inácio P. R. M, Freire M. M. CloudSim Plus: A cloud computing simulation framework pursuing software engineering principles for improved modularity, extensibility and correctness. In: Proceedings of IFIP/IEEE Symposium on Integrated Network and Service Management (IM):400-406; 2017.
20. Lantz B, Heller B, McKeown N. A Network in a Laptop: Rapid Prototyping for Software-defined Networks. In: Proceedings of the 9th ACM SIGCOMM Workshop on Hot Topics in Networks:19:1–19:6; 2010; New York, NY, USA.
21. Riggio R, Khan S. N, Subramanya T, Yahia I. G. B, Lopez D. LightMANO: Converging NFV and SDN at the edges of the network. In: Proceedings of IEEE/IFIP Network Operations and Management Symposium:1-9; 2018.
22. OpenStack Tacker. OpenStack service for NFV Orchestration. <https://docs.openstack.org/tacker/latest/>; 2019; (accessed 05 July 2019).

23. Son J, Buyya R. Latency-aware Virtualized Network Function provisioning for distributed edge clouds. *Journal of Systems and Software*. 2019;152:24–31.
24. Toosi A. N, Son J, Chi Q, Buyya R. ElasticSFC: Auto-scaling techniques for elastic service function chaining in network functions virtualization-based clouds. *Journal of Systems and Software*. 2019;152:108–119.
25. Pelley S, Meisner D, Wenisch T. F, VanGilder J. W. Understanding and abstracting total data center power. In: *Proceedings of the Workshop on Energy-Efficient Design (WEED 2009)*:1-6; 2009; Austin, Texas, USA.
26. Wang X, Yao Y, Wang X, Lu K, Cao Q. CARPO: Correlation-aware power optimization in data center networks. In: *Proceedings of the IEEE INFOCOM*:1125-1133; 2012.
27. Ersoz D, Yousif M. S, Das C. R. Characterizing network traffic in a cluster-based, multi-tier data center. In: *Proceedings of the 27th International Conference on Distributed Computing Systems*:59–59; 2007.
28. Al-Fares M, Loukissas A, Vahdat A. A Scalable, Commodity Data Center Network Architecture. In: *Proceedings of the ACM SIGCOMM Conference on Data Communication*:63–74; 2008; New York, NY, USA.
29. McKeown N, Anderson T, Balakrishnan H, et al. OpenFlow: enabling innovation in campus networks. *ACM SIGCOMM Computer Communication Review*. 2008;38(2):69–74.
30. He T, Toosi A. N, Buyya R. Performance evaluation of live virtual machine migration in SDN-enabled cloud data centers. *Journal of Parallel and Distributed Computing*. 2019;131:55–68.
31. Aroca J. A, Chatzipapas A, Anta A. F, Mancuso V. A measurement-based characterization of the energy consumption in data center servers. *IEEE Journal on Selected Areas in Communications*. 2015;33(12):2863–2877.
32. Louis B, Mitra K, Saguna S, Åhlund C. CloudSimDisk: Energy-Aware Storage Simulation in CloudSim. In: *Proceedings of IEEE/ACM 8th International Conference on Utility and Cloud Computing (UCC)*:11-15; 2015.
33. Ilager S, Ramamohanarao K, Buyya R. ETAS: Energy and thermal-aware dynamic virtual machine consolidation in cloud data center with proactive hotspot mitigation. *Concurrency and Computation: Practice and Experience*. 2019;31(17):1-15.
34. Pedram M, Hwang I. Power and performance modeling in a virtualized server system. In: *Proceedings of 39th International Conference on Parallel Processing Workshops*:520–526; 2010.
35. Kaup F, Melnikowitsch S, Hausheer D. Measuring and modeling the power consumption of openflow switches. In: *Proceedings of IEEE 10th International Conference on Network and Service Management (CNSM) and Workshop*:181–186; 2014.

How to cite this article: J. Son, TZ. He and R. Buyya (2019), CloudSimSDN-NFV: Modeling and Simulation of Network Function Virtualization and Service Function Chaining in Edge Computing Environments , *Journal Title.*, 2019;00:1–6.



Minerva Access is the Institutional Repository of The University of Melbourne

Author/s:

Son, J;He, T;Buyya, R

Title:

CloudSimSDN-NFV: Modeling and simulation of network function virtualization and service function chaining in edge computing environments

Date:

2019-12

Citation:

Son, J., He, T. & Buyya, R. (2019). CloudSimSDN-NFV: Modeling and simulation of network function virtualization and service function chaining in edge computing environments. SOFTWARE-PRACTICE & EXPERIENCE, 49 (12), pp.1748-1764. <https://doi.org/10.1002/spe.2755>.

Persistent Link:

<http://hdl.handle.net/11343/286531>