

Consolidating Virtual Machines by using Ant Colony System algorithm for Green Cloud Computing

by

Md. Kaviul Hossain
17101192

Mutasimur Rahman
17101445

Azrin Hossain
17101398

Samin Yeaser Rahman
17101075

A thesis submitted to the Department of Computer Science and Engineering
in partial fulfillment of the requirements for the degree of
B.Sc. in Computer Science

Department of Computer Science and Engineering
Brac University
October 2020

© 2020. Brac University
All rights reserved.

Declaration

It is hereby declared that

1. The thesis submitted is my/our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

Student's Full Name & Signature:

Md. Kaviul Hossain

Md. Kaviul Hossain
17101192

Mutasimur Rahman

Mutasimur Rahman
17101445

Azrin Hossain

Azrin Hossain
17101398

Samin Yeaser Rahman

Samin Yeaser Rahman
17101075

Approval

The thesis titled “Consolidating Virtual Machines by using Ant Colony System algorithm for Green Cloud Computing” submitted by

1. Md. Kaviul Hossain (17101192)
2. Mutasimur Rahman (17101445)
3. Azrin Hossain (17101398)
4. Samin Yeaser Rahman (17101075)

Of Summer, 2020 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science on October 5, 2020.

Examining Committee:

Supervisor:
(Member)



Dr. Md. Motaharul Islam
Professor
Computer Science & Engineering
United International University

Head of Department:
(Chair)

Prof. Mahbub Majumdar
Chairperson
Dept. of Computer Science & Engineering
Brac University

Mahbubul Alam Majumdar
Professor and Dean
Department of Computer Science and Engineering
Brac University

Abstract

Cloud computing is an internet based computing system enabling convenient, flexible, on-demand network access to a shared pool of configurable computing resources (e.g. networks, servers, storage, applications, services) that can be rapidly provisioned and released with minimal management effort or service provider interactions. Although cloud is a virtual and non-physical network-based concept, its implementation includes physical and wired connectivity among servers, hard-drives and other network devices resulting in large number of data centers. As such power and energy consumptions increase total carbon footprint which is very hazardous for the environment. Green Cloud Computing is a concept of consuming less physical power and saving energy by making the infrastructure more virtual. With the help of the Ant Colony System (ACS) algorithm, it is possible to consolidate the Virtual Machines and make Cloud Computing more virtual than it already is. This algorithm describes a mechanism of indirect coordination, through the environment, between agents or actions e.g. communication about paths between ants with the help of pheromones. Using this technique, the Virtual Machines (VMs) that constitute Cloud Computing can be amalgamated and consolidated. Thus, Green Cloud Computing can be easily globalized.

Keywords: Virtual machine, hyper-visor, state-transition rule, physical machine, meta-heuristic, virtualization.

Dedication

Dedicated to every single soul who works hard to bring a change to the world of science.

Acknowledgement

Firstly, all praise to the Great Allah for whom our thesis have been completed without any major interruption.

Secondly, to our advisor Dr. Md. Motaharul Islam sir for his kind support and advice in our work. He helped us whenever we needed help.

Thirdly, to Marco Dorigo, researcher and inventor of the Ant Colony System algorithm, who helped us in our research by providing necessary resources.

And finally to our parents without their throughout support it may not be possible. With their kind support and prayer we are now on the verge of our graduation.

Table of Contents

Declaration	i
Approval	ii
Abstract	iii
Dedication	iv
Acknowledgment	v
Table of Contents	vi
List of Figures	viii
List of Tables	ix
Nomenclature	xi
1 Introduction	1
1.1 Motivation	1
1.2 Major Contributions	1
1.3 Thesis Overview	2
2 Literature Review	3
3 Background Study	5
3.1 Cloud Computing	5
3.1.1 Salient Features of Cloud Computing	5
3.1.2 Cloud Computing Framework Architecture	6
3.1.3 Cloud Modelling	7
3.1.4 Inter-Networking of Clouds	8
3.1.5 Green Cloud Computing	10
3.2 Virtualization	11
3.2.1 The Process of Virtualization	11
3.2.2 Virtual Machine allocation in Cloud Computing	12
3.2.3 Virtual Machine Migration	13
3.3 Ant Colony System Algorithm	14
3.3.1 ACS Algorithm & its Essential Features	15
3.3.2 OEMACS Algorithm	22
3.3.3 Collaboration of OEMACS & Cloud Computing	25

3.4	Similar Algorithms	26
4	Proposed Model	28
4.1	Active & Idle Virtual Machine Migration Algorithm	29
4.2	Working Model	29
4.3	Salient features of AIVMM	33
4.3.1	A multi-objective approach towards VMM	33
4.3.2	Service Level Agreement (SLA)	33
4.3.3	Resource Utilization	34
4.3.4	Energy Utilization	34
4.4	Limitations of AIVMM	35
5	Result & Analysis	36
5.1	Result	36
5.2	Analysis	36
5.2.1	Data Analysis of OEMACS & ACS Algorithm	36
5.2.2	Data Analysis of AIVMM Algorithm	39
6	Conclusion & Future Work	42
6.1	Conclusion	42
6.2	Future Work	42
	Bibliography	43

List of Figures

3.1	Cloud Computing	5
3.2	Features of Cloud Computing	6
3.3	Application Layer of Cloud Computing	7
3.4	Cloud Modelling	8
3.5	Inter-Networking of Cloud	9
3.6	Energy Consumption by Cloud Data Centers	10
3.7	Green Cloud Computing	11
3.8	Virtualization	12
3.9	Scheduling Policy	13
3.10	Virtual Machine Placement Problem	15
3.11	Propagation of real agents (ant)	16
3.12	Propagation of artificial ants	17
3.13	Flowchart of ACS Algorithm	20
3.14	The system architecture of OEMACS	22
3.15	Order Exchange & Migration	25
5.1	Average Time Complexity of OEMACS with Maximum Iteration 150	37
5.2	Average Time Complexity of OEMACS with Maximum Iteration 500	37
5.3	Average Time Complexity of OEMACS with Maximum Iteration 1000	38
5.4	Average Time Complexity of OEMACS with Maximum Iteration 2000	38
5.5	Time to Simulation Instantiation	40
5.6	Memory to Simulation Instantiation	40
5.7	OEMACS Vs (OEMACS+AIVMM)	41

List of Tables

5.1	Dataset for ACS algorithm Simulation	36
5.2	Time to Simulation Instantiation	39
5.3	Memory to Simulation Instantiation	39
5.4	OEMACS Migration	41

Nomenclature

The next list describes several symbols & abbreviation that will be later used within the body of the document

ACS Ant Colony System

AIVMM Active Idle Virtual Machine Migration

AS Ant System

BFD Best Fit Decreasing

CDN Content Delivery Network

CE Cloud Exchange

CO Carbon Monoxide

CO₂ Carbon-di-oxide

COP Combinatorial Optimization Problem

CPU Central Processing Unit

EC Evolutionary Computation

GA Global Agent

HACOPSO Hybrid Ant Colony Optimization with Particle Swarm Optimization

IaaS Infrastructure as a Service

IDE Integrated Development Environment

LA Local Agent

MACO Multi-objective Ant Colony Optimization

MIMPS Migration of Idle Machines via Parking Server

NIC Network Interface

OEM Order Exchange and Migration

OEMACS Order Exchange Migration Ant Colony System

PaaS Platform as a Service

PM Physical Machine
PSO Particle Swarm Optimization
QoS Quality of Service
RAM Random-access memory
SaaS Software as a Service
SLA Service Level Agreement
SLAV Service Level Agreement Violation
TSP Travelling Salesman Problem
VM Virtual Machine
VMC Virtual Machine Consolidation
VMM Virtual Machine Monitor
VMP Virtual Machine Placement
VMs Virtual Machines
XML eXtensible Markup Language

Chapter 1

Introduction

1.1 Motivation

Cloud Computing is the future of online storage and processing. The most important feature of today's storage system and remote accessibility is its ubiquitous availability and multi-tenancy, which grows more valiantly with the help of cloud computing. The salient features and gradual development of cloud computing was our main motivation to work in this field. Communication, storage and accessibility has become piece of cake because of the bulky characteristics of Cloud Computing. Big data management has been a notable success of this field. Although there remain some drawbacks in the system, the problem with high power and energy consumption by the idle virtual machines and emission of harmful gases like CO_2 , CO, Methane etc. in a data center has motivated us to work on it and leave noticeable changes in that aspect. Our motive is to design an algorithm that will not only work perfectly with the current one but also solve the high-power consumption problem with remarkable changes. From our background study of the system, we learnt that many researchers have worked with the problem at different times and they came up with both single and multi-objective algorithms to fix the problem. Thus, we aimed higher and tried to introduce a self-explanatory system or algorithm called Active and Idle Virtual Machine Migration algorithm that not only solves the existing glitch but also makes the process more efficient and consolidated. We have discussed our contributions elaborately in the next section (1.2).

1.2 Major Contributions

Most researchers have introduced different algorithms to address the problem of excess energy consumption of idle virtual machines. They have used algorithms like multi-objective ant colony optimization (MACO), hybrid ant colony optimization with particle swarm optimization (HACOPSO) and much more to fix the problem. [1] But at a certain point, the algorithms fail to solve the problem with all the objectives intact. In order to sort that problem and address the existing drawbacks of the algorithms, we came up with such an algorithm that can consolidate the cloud architecture by making virtual machines consume energy in an efficient way and also keep the integrity of the system intact and unchanged. To do such, we have maintained integrity of the cloud architecture and consolidated the virtual machine placement problem in collaboration of a meta-heuristic algorithm called the Order

Exchange Migration Ant Colony System (OEMACS). We have also ensured Green Cloud Computing by reducing high thermal throttle in physical machines like servers that host the Virtual Machines in a data center. We successfully migrated idle or zombie VMs from one server to another containing maximum idle VMs, which has ensured no unwanted lag of power supply for the actively working VMs.

1.3 Thesis Overview

The subsequent sections of the thesis have been organized as follows. Chapter 2 is the literature review which contains related existing approaches along with the problem statement relevant to our proposed model. Chapter 3 includes all the background information related to our work and how we are using the resources to get the desired output, along with a review of all the existing similar models that researchers came up with at different times. Section 3.1 contains the introduction to Cloud Computing along with its salient features and architecture. Section 3.2 is the section discussing the process of virtualization and its collaboration with Clouds. Section 3.3 is the elaborated narration of the Ant Colony System Algorithm used to design our proposed model. In chapter 4, we have described the proposed model of our works along with relevant graphs and figures. It includes the overall models relevant to our solution algorithm. The predicted results and relevant discussions are showed in chapter 5. Lastly, the summary of the report and conclusion as well as some future work-plan is done in chapter 6. In future work part, we will talk about our future ambition and working plan to upgrade our proposed model for better performance.

Chapter 2

Literature Review

Cloud Computing has made our lives easy and simple as we can access our important files, documents, folders and all other on the go resources, anywhere and everywhere we go. But the system itself is not that simple as it seems. Cloud Computing is a complex computing technique with a sophisticated architecture and vast service domain. The whole architecture can be divided into two sub-domains: physical and virtual. A cloud computing system is comprised of both physical machines (servers) and virtual wares (VM-wares). With the help of an intermediate application named Hypervisor, the mainframe of a server is converted into 'N' virtual machines. A server has four basic components: Processing unit, Memory unit (RAM), Storage unit (Hard Disc) and NIC (Network Interface). Hypervisor virtualizes these four components of a physical machine and divides them into 'N' virtual machines. The value of n for every physical machine depends on the configuration of the four basic components a server is composed of.

Cloud computing is a modern concept that deals with new computing paradigm. It supports new and modified features of big data management, data storage, data processing as well as data multi-tenancy. A cloud computing architecture includes both physical and virtual components like data centers, servers, virtual machines, migrating algorithms and much more. There are mostly three service models that cloud computing has been providing from the very beginning. They are: Infrastructure as a Service (IaaS), Platform as a Service (PaaS) and Software as a Service (SaaS). Along with these services, the cloud paradigm is deployed in three ways. Public, private and hybrid. Again, by mixing two or more methods, cloud can be deployed in some more ways. Community cloud is one such example. Using its pay-per-use feature as a business model, cloud computing converts the investment risk of under or over provisioning to the cloud computing service providing companies. As such, numerous public and private Infrastructure, Platform and Software as a Service cloud vendors such as Microsoft, Google and Amazon, tend to operate big-scale data centers all over the globe. Moreover, in order to cope up with the rapid increment in cloud infrastructure and service demand, these vendors or cloud service providers have made a very significant increase in the power, size and energy consumption of the data centers. Thus, the cost to maintain and handle the increment in power and infrastructure has gone up to a great extent. Not only the power consumption but also the equipment cooling costs has risen a lot. Microsoft has taken the initiative on implementing under water data centers so that the cooling technique becomes easy and efficient. Previous researches on power consumption by

cloud computing data centers have shown that in the year 2010 data centers around the globe have consumed nearly two hundred and one Kwh of electricity. In the recent times, almost all the cloud data servers utilize between 11-50 % of the full resources. As such, trying to place the idle virtual machines of a low utilized server on other active servers will reduce the power consumption problem to a great extent.

In recent years, some attempts have been made to reduce the energy consumption of data centers. One effective and commonly use method is consolidation. Server consolidation is a technique of reducing the energy cost of the data center by migrating active but idle virtual machines to fewer servers. All the active but not working virtual wares are migrated to a single server and put into low power mode while the actively working virtual machines are migrated to a single server. As a server with ‘N’ fully loaded virtual machines are taking the full load whereas the idle machines are hibernating in a completely different server. In the recent times, many researchers and cloud specialist came up with different efficient algorithms and techniques to consolidate a server. They mainly focused on the load balancing problem by migrating virtual machines from one server to other. The primary and foremost concentration of these consolidation ideas is the continuous migration and reallocation of idle VMs in the dynamic infrastructure.

In our paper, we are focusing on executing the server consolidation method by combining Order Exchange and Migration (OEM) paradigm with a complex but efficient algorithm called the Ant Colony System (ACS) algorithm. Ant Colony System algorithm or the ACS algorithm is a versatile, robust, heuristic and population-based approach towards solving the Travelling Salesman Problem (TSP). In this algorithm, ‘m’ ants (algorithm agents) focus on determining the shortest path from their source i to destination j by undergoing s iterations in time t . In real life, ants use pheromones to mark their paths between their food source and homes. The level of pheromones on their paths act as their stigma of influencing more ants to follow that path. Using this very unique techniques of ants, Mark Dorigo (1961) developed an algorithm of shortest path trace in machines called the ‘Ant Colony System’ algorithm.

Chapter 3

Background Study

3.1 Cloud Computing

Cloud computing is a modern-day computing paradigm enabling convenient, flexible, standby online access to a multi tenant shared storage of many efficient and useful computing resources which can be quickly and effectively managed and released with top-notch management service and handles with vendors' collaborations. It is the process of storing, sharing and controlling data and files in a remotely manner where a user can share their data with a wide pool of other users. Although cloud is a virtual and non-physical network-based concept, its implementation includes physical and wired connectivity among servers, hard-drives and other network devices resulting in large number of data centers.

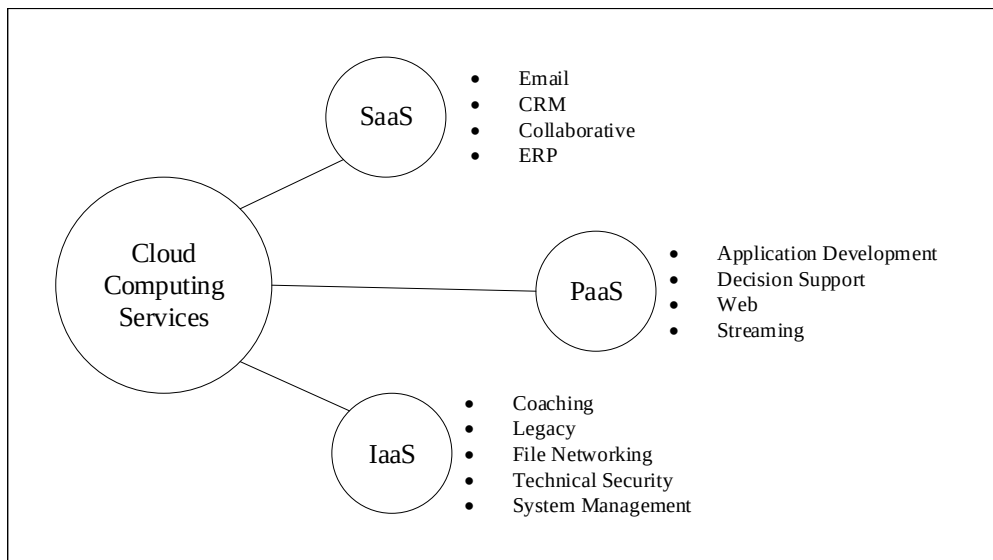


Figure 3.1: Cloud Computing

3.1.1 Salient Features of Cloud Computing

Salient features of cloud computing include on-demand service, resiliency, ubiquitous access, multi-tenancy, self-serviced and pay as you need properties. Cloud is an enormous entity hosting more than ever seen before data that can be accessed from

anywhere on earth through internet. Anyone anywhere at anytime can access the cloud and store or preserve any data needed. A user can not only store data in a cloud but also share the data with a pool of other users which makes cloud computing a multi-tenant concept.

Whenever a server faces power shutdown due to power failure or regular maintenance, the data is stored and migrated to another server. As such, a customer or user can access the data anytime they want even if a server is down. Cloud computing is resilient technique of data management. One can easily use this platform and become an expert in it. It is a ubiquitous computing model that allows a person to access and use it from any corner. Any major failure of power or connectivity does not harm the overall architecture and its functionality.

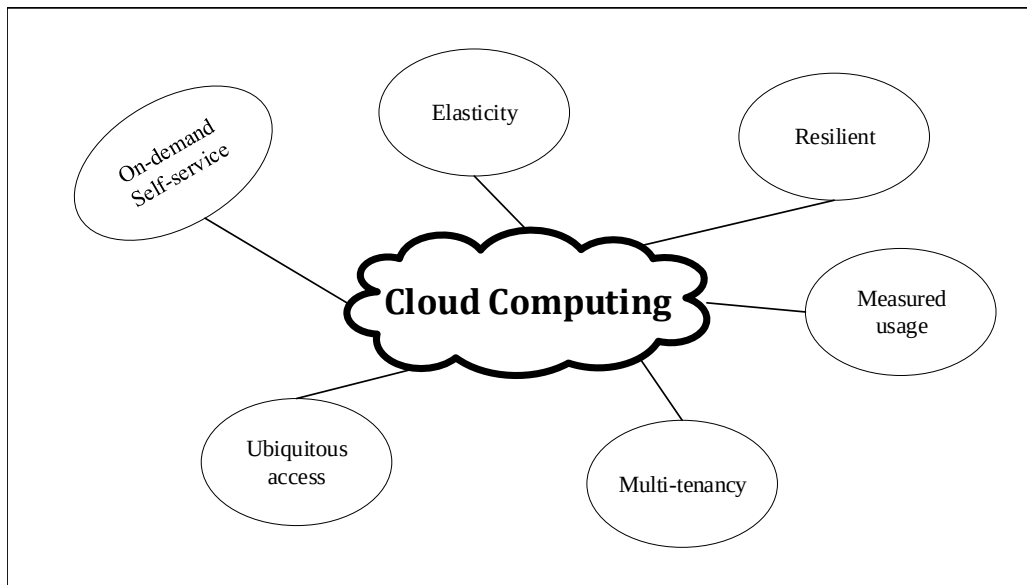


Figure 3.2: Features of Cloud Computing

3.1.2 Cloud Computing Framework Architecture

The resources of physical cloud and center middle-ware abilities together provide the reason for conveying IaaS and PaaS. SaaS capabilities are provided by the user-level middle-ware. The SaaS/PaaS service providers are different from the IaaS service providers, usually third-party service providers develop PaaS/SaaS services.

Cloud applications: Cloud layer consists of applications that can be accessed by the end-users. End-users can use the SaaS applications over the internet which is provided by the SaaS providers. In this layer, the users can send their possess application as well.

User-level middle-ware: Different software frameworks, for example Ruby on Rails, Asynchronous JavaScript and XML (AJAX) are incorporated in this layer. These frameworks help developers to build a project efficiently. This layer provides platforms for cloud software development, modeling, documentation and other composition tools as well. Additionally, frameworks from this layer help the application from the upper level.

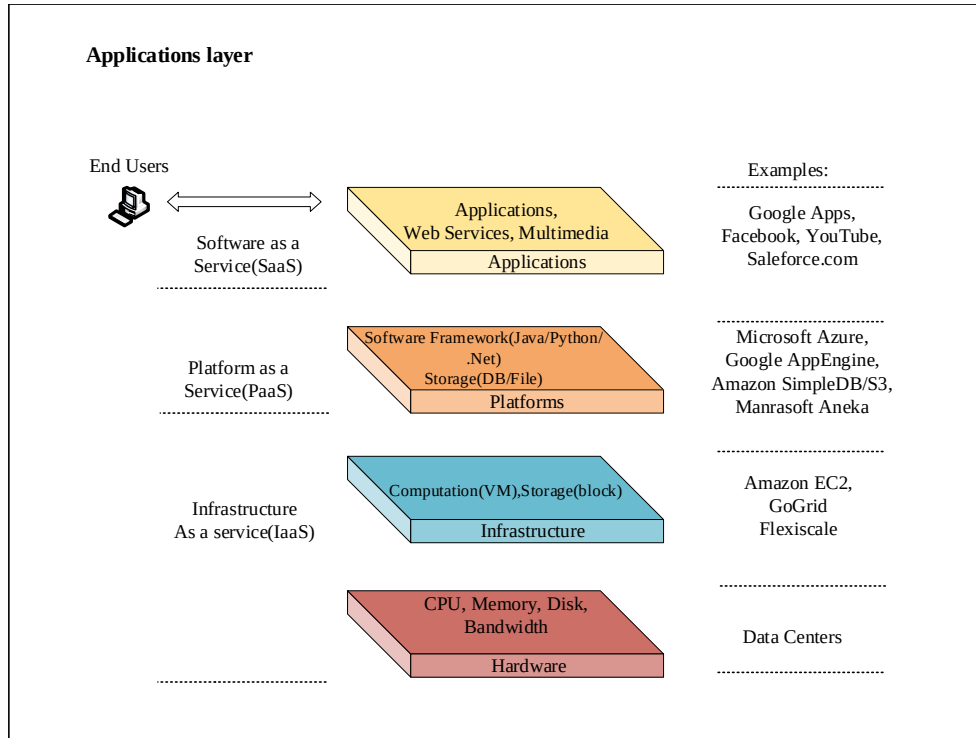


Figure 3.3: Application Layer of Cloud Computing

Core middle-ware: For facilitating the user-level services, the core middle-ware provides an appropriate run-time environment for applications. The core middle-ware services include remote management, execution management and monitoring, storage system, security services. Few of the references are Manjrasoft Aneka, GoGrid, Amazon EC2 and S3. Manjrasoft Aneka provides the customer with a floor for application development purposes, that are hosted in cloud. GoGrid or Amazon EC2 and S3 offers the customer with virtual hardware and storage on top of which they can build their infrastructure (IaaS).

System level: This is the bottom layer of the cloud architecture that focuses on Infrastructure-as-a-Service (IaaS). It provides access to computing resources that includes storage servers and application servers. These servers are handled by top-level virtualization.

3.1.3 Cloud Modelling

CloudSim is a framework for modeling and simulation of cloud computing infrastructure. According to CloudSim, entity is an instance of a component. IaaS can be simulated by increasing data center entity. Data center entity controls host entities. The hosts are allocated to one or more VMs. The VM allocation policy assigns the hosts to VMs. Comparatively, one or more application services can be allocated within single VM instance. A data center can control few hosts. Host works as a physical server and a component of CloudSim.

VM allocation generates VM instances on hosts that match software environment, specifications and characteristics of the SaaS provider. CloudSim supports the improvement of custom application service models and the clients have to expand the

core cloud-let object for executing their application service[14]. CloudSim offers customizable policies for provisioning host and resources to virtual machines on a first-come-first-serve service, hardware requirements. Although, policies used by public cloud providers are not available publicly.

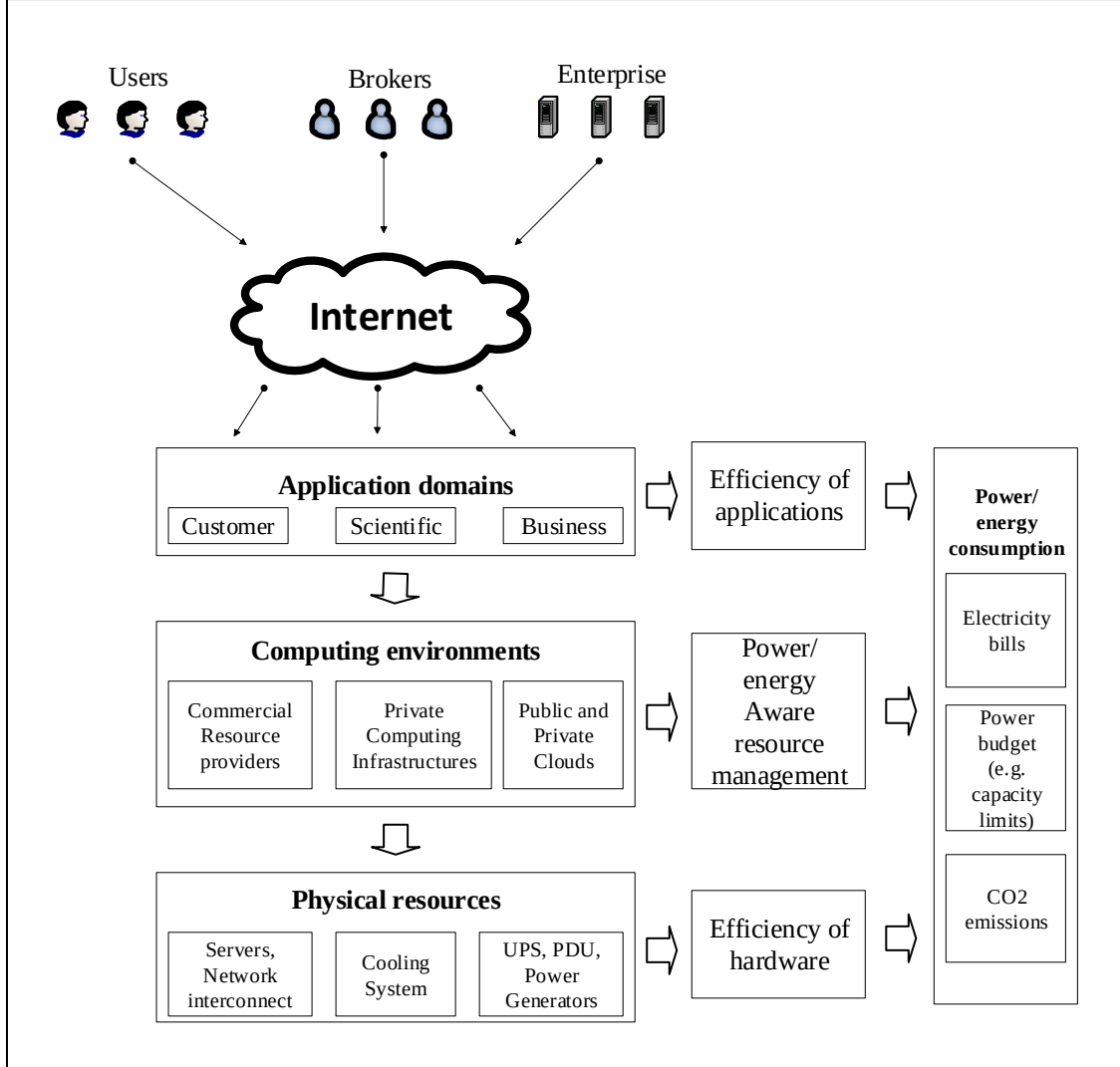


Figure 3.4: Cloud Modelling

Host allocation policy does the provision of processing cores to VMs for each host component. The policy comprehends several hardware requirements such as memory storage, CPU share, number of CPU cores. CloudSim supports the idea of allocating particular CPU cores to particular VMs, assigning the capacity of a core among VMs dynamically or on demand. These ideas are known as space-shared policy and time-shared policy. Each host component represents a VM scheduler component. The scheduler component can be executed either by the space-shared policy or the time-shared policy.

3.1.4 Inter-Networking of Clouds

To ideally serve client needs around the world, cloud service suppliers have a few centers at distinctive location over the internet. Federated cloud provides remarkable

execution and economic benefits by enhancing the capability of SaaS providers and present better service by optimizing the service. Cloud providers do not need to set up a new data center in every location since every member is allowed to obtain additional resources from federation dynamically.

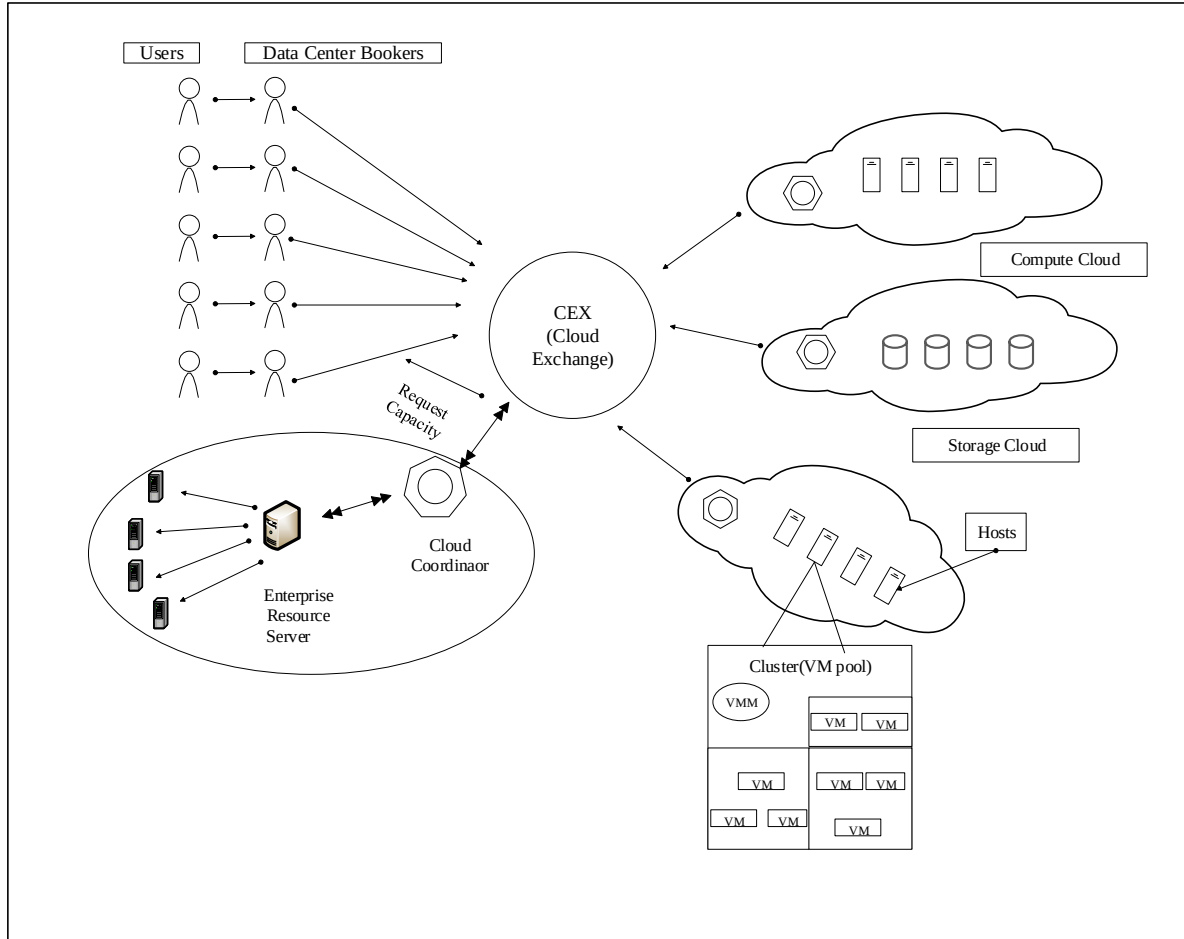


Figure 3.5: Inter-Networking of Cloud

Cloud coordinator is one of the major elements of clouds and their federated network. Cloud coordinator is responsible for exporting cloud services to the federation by implementing basic functionalities such as scheduling, virtualization, dynamic monitoring, application composition etc. Cloud Exchange (CEX) facilitates the cloud brokers acting on behalf of the SaaS providers by identifying appropriate cloud service[16]. Various social networks and content delivery networks (CDNs) may benefit from the federated cloud service architecture[19]. Social networking sites like Facebook, MySpace are using multi-tiered web application. These sites are providing dynamic services to lots of end-users. All component will run on a different virtual machine that can be hosted in data centers maintained by different cloud service providers. Every developer can choose their cloud service provider according to their plug-in. As a result, a social networking site is structured by a large number of diverse services that may be hosted by lots of data centers around the world.

3.1.5 Green Cloud Computing

The process of ensuring low power consumption and high resource utilization along with proper and regular outcome is basically green computing. The concept of green cloud computing is a part of green computing as servers and physical hardware used in the process emits CO₂ and other carbon-based gases that are highly hazardous for the environment. According to Bioscience Biotechnology Research Communication (BBRC), consumption of power and energy by server in cloud data centers are increasing at an alarming rate. Fig.(3.6) shows the scenario in a graphical form.

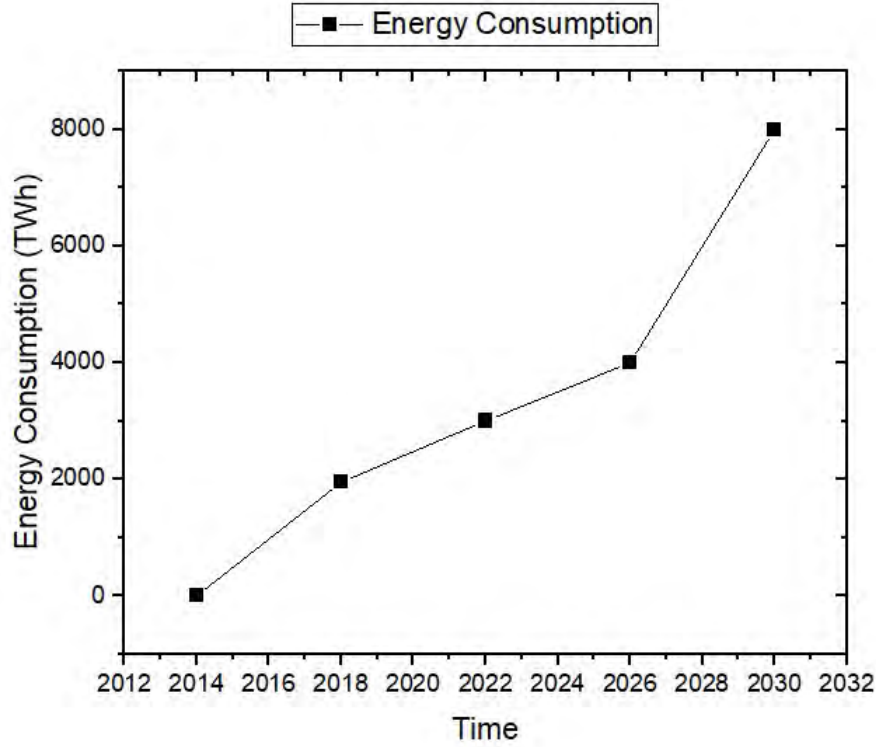


Figure 3.6: Energy Consumption by Cloud Data Centers

In order to ensure a green cloud computing environment, we have to stop the high-power consumption of idle virtual machines in the cloud servers. Whenever a VM is assigned a task, it takes memory and processing powers from the under-laying servers. After the task is completed, the VM becomes idle but it does not stop consuming that resources completely. Rather it keeps using the resources mostly memory to stay idle. As such, power is consumed which causes more greenhouse gas emission.

In order to stop excessive green-house gas emission and ensure green cloud computing, we introduced the Active & Idle Virtual Machine Migration algorithm that helps virtual machines to utilize resources to a great extent and makes idle virtual machines migrate to a server with highest number of idle VMs in it.

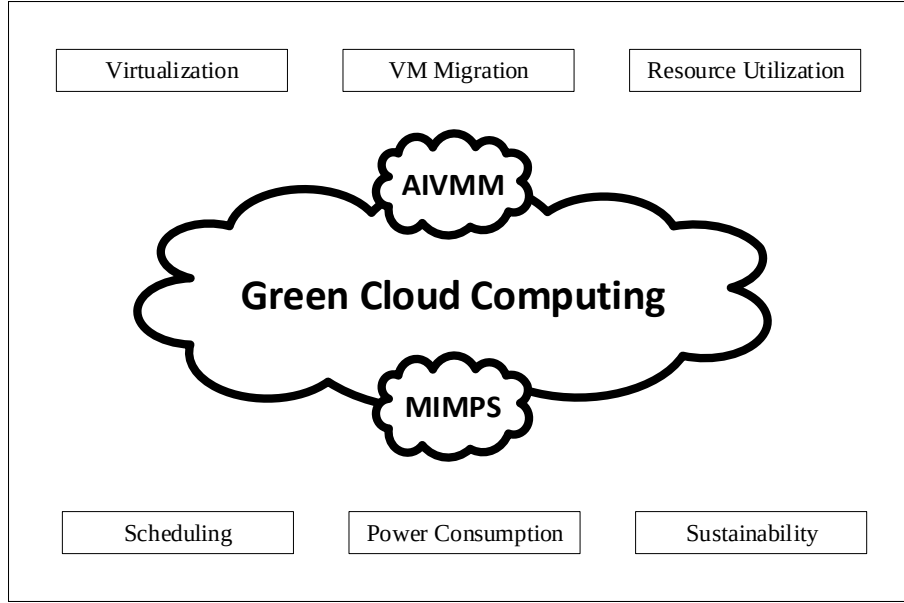


Figure 3.7: Green Cloud Computing

3.2 Virtualization

3.2.1 The Process of Virtualization

Virtualization is the process of virtualizing an operating system or an entity like software, hardware, kernel etc. with the help of an intermediary application known as ‘hypervisor’ so that the resources of a physical machine can be properly utilized. Virtualization is a modern technique which is very efficient and progressive. With the help of virtualization, physical resources like RAM, CPU and hard-drive storage can be properly utilized[22][23].

The process of virtualization is of two types- full virtualization and para-virtualization. The under-laying physical machine with the main operating system is called the host machine. With the help of a software or application called the hypervisor, the host operating system is virtualized and so are the physical resources. As such, more than one virtual machines are created on top of the physical host and each of the virtual machines are assigned processor, memory and storage.

There are four major components that a machine needs to operate properly. They are processor, memory, storage and network bandwidth. While installing a virtual machine over a physical machine, these four components are allocated to the VMs. A virtual machine can perform all the tasks that a physical machine can. The main difference between these two types of machine is the virtualized property itself.

The process of virtualization in which the hypervisor uses a primitive and unchanged version of a guest OS is called a full virtualization. The guest machine connects with the Central Processing Unit of the physical machine through a channel that the hypervisor creates for addressing purpose. Through this channel the guest OS can directly communicate with the CPU of the host. This type of virtualization is the fastest. On the other hand, the paravirtualization process uses a changed and edited version of the guest OS. In paravirtualization, the virtual machine communicates with the host Central Processing Unit through the hypervisor. The hypervisor

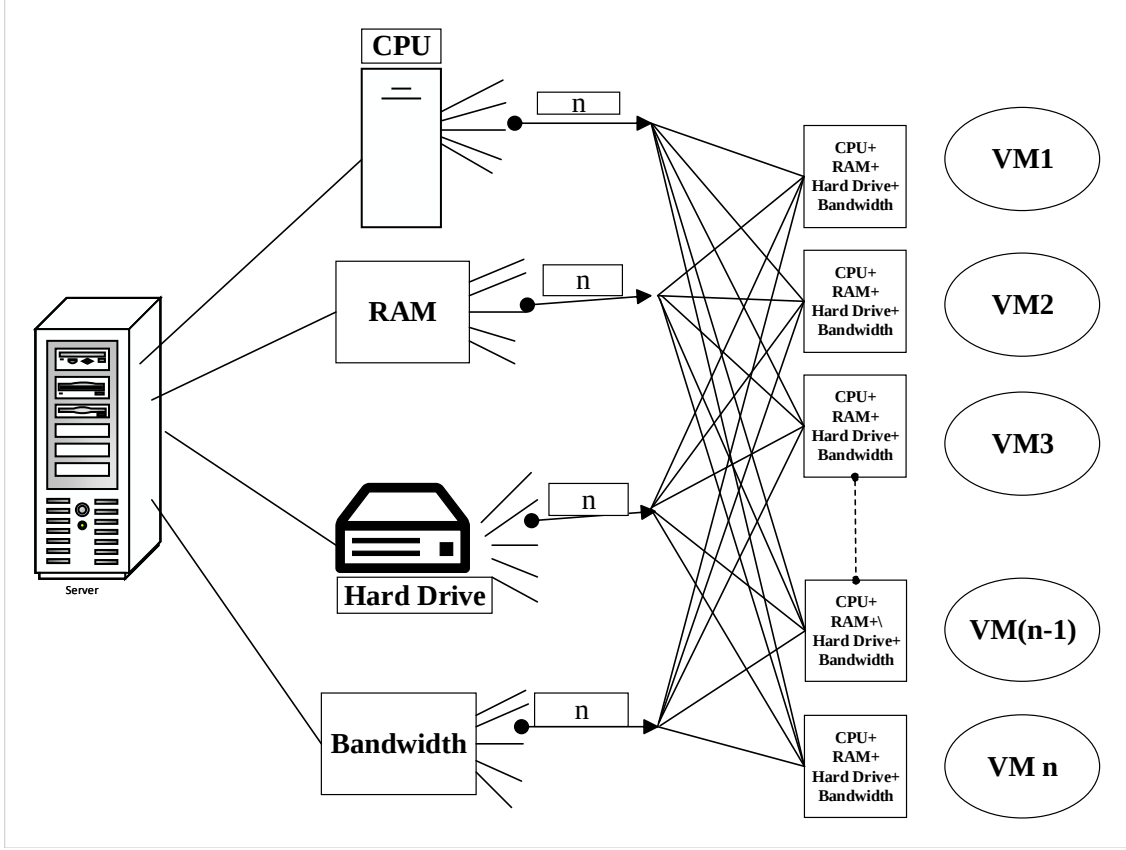


Figure 3.8: Virtualization

channels the unchanged calls from the VM to the host and other possible interfaces.

3.2.2 Virtual Machine allocation in Cloud Computing

The infrastructure of cloud computing and that of grid computing is different because of virtualization tools. Additionally, cloud has an extra virtualization layer. VMs are isolated yet they share processing cores and system bus. Usually, each VM use the processing power that is available within the host. However, VM can demand more processing power than the availability. To resolve this issue, CloudSim offer VM allocation in two different levels: host and VM level. At first level, the processing power of each core of each VM is determined. Then, at VM level, the VM distributes the available power to individual application services. Here, we will assume a task unit as an individual application service within the VM.

With the help of a software named CloudSim, it has become extremely easy to allocate virtual machines in the Cloud architecture. There are two provisioning policies maintained by CloudSim[16]. They are- time-shared and space-shared policies. CloudSim applies these policies in all the levels. In time-shared policy, a VM completes the tasks assigned to it with respect to time. It complete one task at a time and moves to the new one. While in space-shared policy, a VM can complete more than one cloudlet (task) simultaneously by sharing the resources. fig 1.8 depicts an allocating scenario in which the space-shared policy has been applied to both virtual machines and cloudlets.

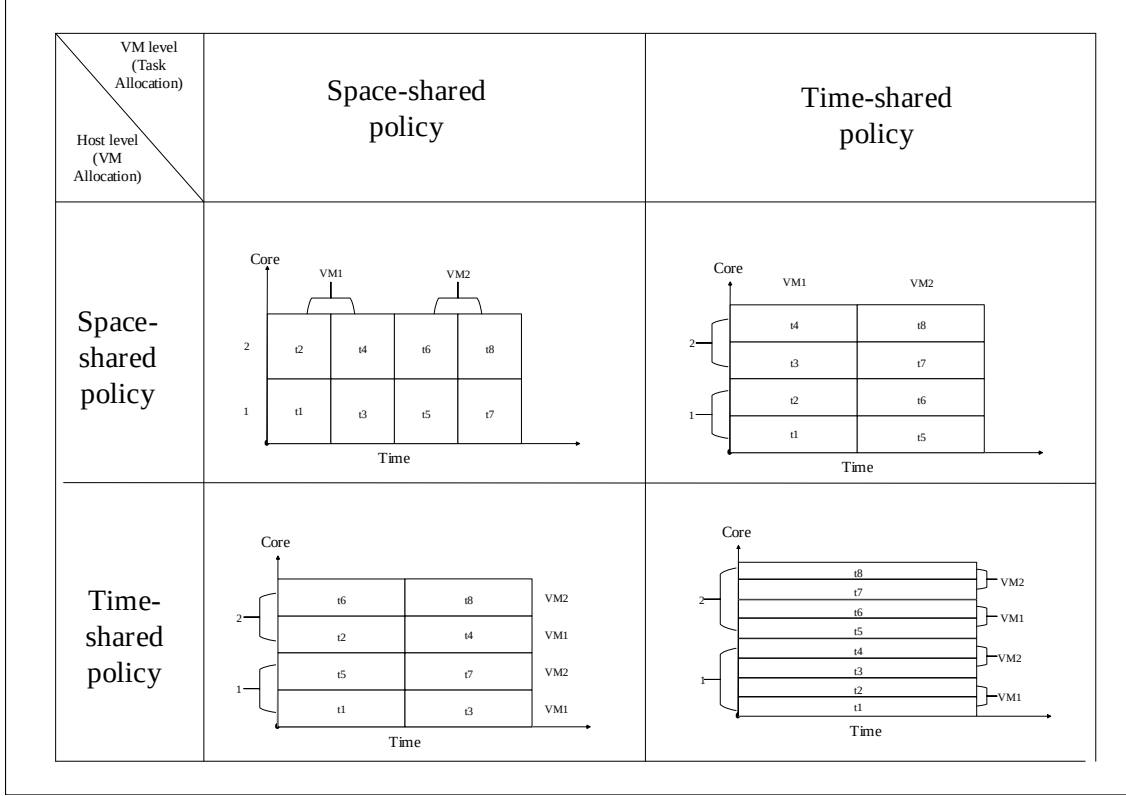


Figure 3.9: Scheduling Policy

3.2.3 Virtual Machine Migration

Virtualization permits dividing one host machine into more than one virtual machines that share the same physical resources and runs simultaneously as well. Due to avoid load balancing and make the physical machine fault tolerant, virtualization allows virtual machine (VM) to migrate from one physical machine to another physical machine. To reduce energy consumption of cloud data centers, virtual machine consolidation is an emerging solution. VM consolidation is a technique that combines various virtual servers to run on a physical server at the same time.

Live Virtual Migration:

Live virtual migration is the movement of a guest machine from one host server to another server without any downtime. It helps in load balancing and also ease VM installation across physical machine in data center. This is effective in several schemes. VMs can be migrated from a failed physical machine to a stable one[18]. Idle VMs can be moved to another physical machine to optimize resource utilization. In case of load balancing, VMs can be migrated across multiple machines.

Live migration can take place in two processes:

- switching the control from source host to destination host
- move the data storage including virtual disks to the destination host.

The two common approaches can be distinguished by the processes mentioned earlier:

1. **Pre-copy:** At first, the memory is shifted and then the execution is shifted. The memory is shifted by a course of iterations. When the migration proceeds, the memory pages starts to get checked and then transferred from the source host to the destination host. The VM runs on the source host while shifting the memory pages. The pages that have been dirtied (changed) during the iteration are re-copied to the destination until the number of dirty pages is lower than the re-copied pages. After that, the VM will stop on the source host and the last memory pages are shifted to the destination. The execution state like CPU state, registers are shifted to the target host. Then the VM will resume on the destination host. In pre-copy, higher number of dirty pages can slow down the performance.
2. **Post-copy:** In post-copy migration process, first the execution system is shifted and followed by that, the memory is shifted. The target virtual machine is halted on the source. After that, a part of the execution state is shifted to the destination first. In the destination, the execution is resumed. If the target machine wants to access any page that has not been moved to the destination yet, then it creates a fault page known as network faults. Network pages fault can bring degrade in performance. In post-copy, the state of VM is assigned over both source and destination host.

Total migration time, down time, preparation, resume these factors are used to examine the performance of live migration.

There is another hybrid method that brings out the best from pre-copy and post-copy method. There are few state information like devices and frequently accessed pages of the VM, is added to the processor state. This approach decreases the number of network fault pages. This method also transfers a set pages in its locality when network fault page is found.

3.3 Ant Colony System Algorithm

In 1992, an Italian researcher named Marco Dorigo proposed a surreal but very realistic algorithm in this doctoral thesis. The algorithm is called the Ant Colony System (ACS) algorithm and it aims to look for optimal paths in a graph. The algorithm searches for paths based on the innate behaviour of ants in which they tend to look for their food sources and once found they follow a pheromone based trace to bring their food from the source back to their nests. In the field of operational research, the ACS algorithm is a NP-hard problem-solving algorithm. Not only that, with the help of this algorithm, Computational Optimization Problems (COP) can be solved by finding the optimal solution from a finite set of solutions. In order to choose the best solution from a given or existing set of solutions, an EC algorithm is inevitable. The Ant Colony System Algorithm is such an EC algorithm.

The main reason of choosing this Ant Colony System Algorithm for our research is that our main target problem i.e. the Virtual Machine Placement problem is a COP as well as an NP-hard problem that requires an EC algorithm. Of all the EC algorithms, the ACS algorithm addresses the problem more efficiently with strong logical inference.

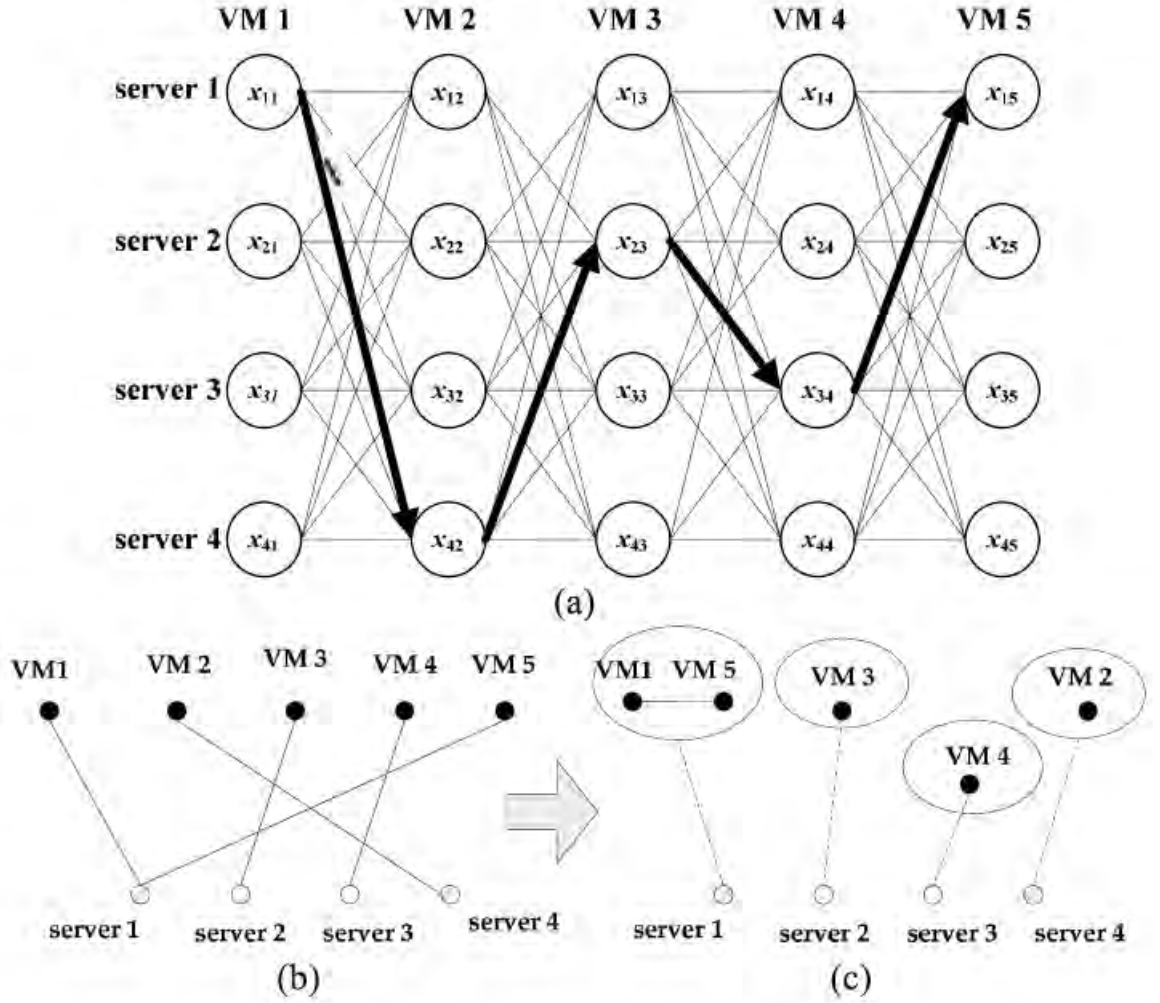


Figure 3.10: Virtual Machine Placement Problem

3.3.1 ACS Algorithm & its Essential Features

Ants are disciplined, industrious, blind insects that work very hard to gather food all summer so that they have enough food stored in the cold seasons. Of all insects and bugs, ants show the most systematic display of unity in performance. Ants are not only stigmergic but also catalytic. In fact, ants are auto catalytic, and they automatically influence their comrades (other ants) to get involved in any altruism thus helping themselves.

In our research, we are using such a system where placement problem of virtual machines in servers will be solved using stigmergic algorithm. The system that we are following in order to consolidate virtual machines and migrate them to active but idle servers, is called the Ant Colony System Algorithm (ACS). Although there are some differences between a real ant and the artificial ants (agents of the algorithm). The advantages that the artificial ants have are:

- They have short-term memory

- They have vision and they can see unlike real ants
- They dwell in an environment in which time is not continuous.

In an Ant Colony System algorithm, an artificial ant acts almost the way a real ant behaves when it moves between two points, its home and remote food source. In order to choose the shortest path between these two points, the ant uses trails of pheromones to compare all the possible paths between them.

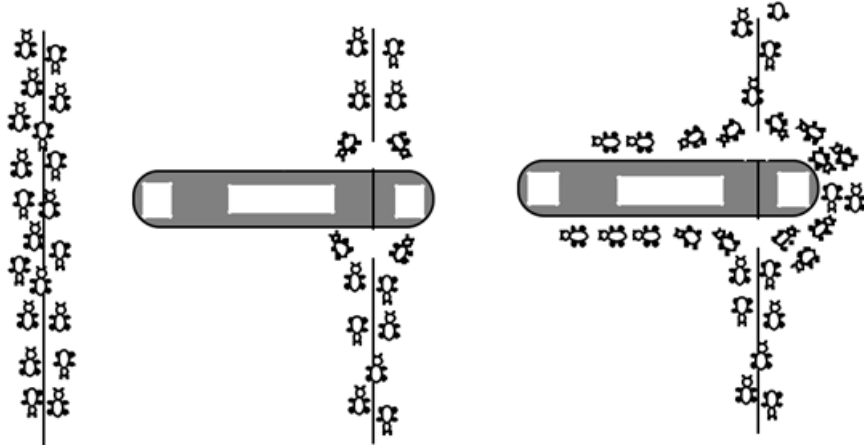


Figure 3.11: Propagation of real agents (ant)

- Here real ants travel along the path that lies between A and E.
- When obstacle appears; agents have the freedom to choose any of the two different paths having equal probability.
- Ants lay the most amount of pheromone on the shortest possible path.

In order to understand the concept of ant system and how ants communicate, travel and collect food, let us take the above fig.(3.11). Let us consider that the distances between point D & H is equal to the distance between point H & B which is 1. Also, let us consider the distance between point B & D via C is equal to 1. Let C be the middle point between D & B fig.(3.12a). Now since time is discrete in our algorithm, we can measure the events happening at different fraction of time by taking $t = 1, 2, 3 \dots t$. In the beginning, let us assume that 30 ants are coming to point B from A and 30 to point D from E at each time unit, such that each ant travels with a velocity of 1 per time unit, and while moving an ant leaves a pheromone trail of intensity 1 (at time t), which tends to evaporate completely and instantaneously in the middle of the successive time interval ($t+1, t+2$).

When $t=0$ there happens to be no trail, as no ant has traveled through that path before but 30 ants are in B and 30 in D. Their choice of choosing path is completely random. Therefore, on an average the ants divide into two groups of 15 and one group moves towards point H while the other moves towards point C. fig.(3.12b). When $t=1$, 30 new real ants start to come from E and A. After they reach point B and D an intensity trail of 15 units has been found. Similarly an intensity trail of 30 units has been found in the path between B and C. The reason for trail being high

on path BC and DC is that, the distance between these three points is less compared to the paths between D, H and B. Since we considered all the ants moving at a same speed, the only variable in this scenario is the distance between the points. As such, by the time ants move from B to H then to point D (and vice versa), ants on the path BC and CD take half the time to complete the route BCD and DCB.

The process continues until all of the ants eventually choose the shortest path.

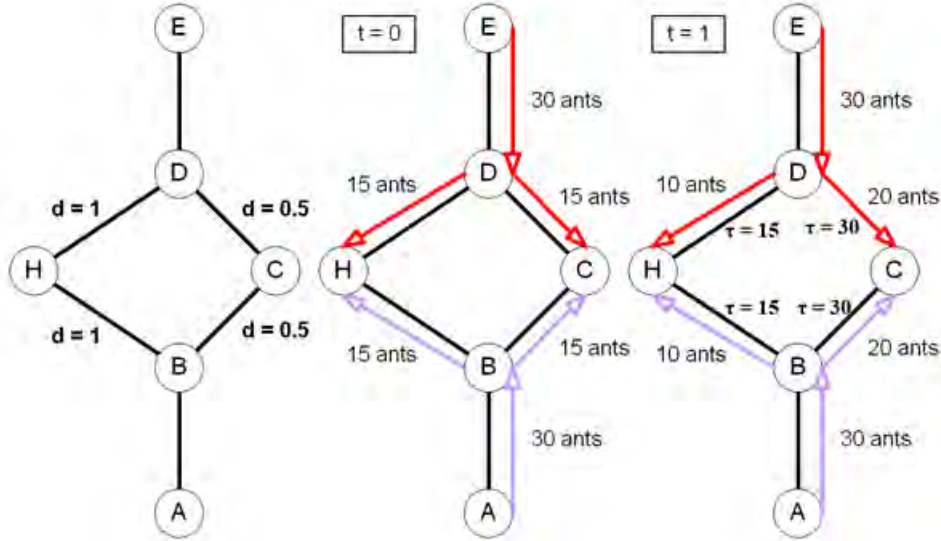


Figure 3.12: Propagation of artificial ants

- First graph depicts the initial graph with the values of distances.
- Second graph shows that when time is equal to zero there remains zero trace on the paths; which means that agents have the choice of choosing to turn either right or left. Here both sides have same probabilities.
- Last graph shows that at time $t=1$ trail is stronger on the short edges, as the ants prefer the short paths while propagation.

The concept is that if at a point an ant has to decide which path to choose, it will always choose the paths with high pheromone trails in it as high amount of pheromone means high probability. Also, high trail means that the path is shorter as ants prefer shorter path. This very trait of following ancestors or activities of preceding companions is called Stigmergy.

Ants as agents: The most common problem solved by using the ACS algorithm is the Travelling Salesman Problem (TSP). In this problem each ant has the following characteristics:

- The ant(agent) decide to travel to the town that has a probability of having shortest distance and most amount of pheromone.
- In order to make a tour legal, visited towns are disabled until a tour is complete (which is controlled by a list called the tabu list[7]).

- When an ant completes a tour, it lays a substance called pheromone trail on each edge (i, j) visited.

The symmetric TSP has a Euclidean based problem space. We use d_{ij} to denote the distance between any two cities in the problem. As such

$$d_{ij} = [(x_i - x_j)^2 + (y_i - y_j)^2]^{1/2} \quad (3.1)$$

Let $\tau_{ij}(t)$ denote the intensity of pheromone trail on edge (i, j) at time t. At the end of each algorithm cycle, the trail intensity between two nodes or cities are updated. Each ant subsequently deposits trail of quantity Q/Lk on every edge (i, j) visited in its individual tour. The sum of all newly deposited trail is denoted by $\Delta\tau_{ij}$. Following trail deposition by all ants, the trail value is updated using the following formula:

$$\tau_{ij}(t + n) = p \times \tau_{ij}(t) + \Delta\tau_{ij} \quad (3.2)$$

Here, p is the rate of trail decay per time interval, τ_{ij} denotes the change in trail intensity and

$$\Delta\tau_{ij} = \sum_{k=1}^m \tau_{ij} \quad (3.3)$$

In the TSP, two factors drive the probabilistic model:

- Visibility, denoted by η_{ij} , equals the quantity $1/d_{ij}$
- Pheromone Trail, denoted by $\tau_{ij}(t)$

These two factors play an essential role in the central probabilistic transition function of the Ant System. In return, the weight of either factor in the transition function is controlled by the variables α and β , respectively. The Probabilistic Transition function used by each ants to decide the next node or city to travel is:

$$P_{ij}^k(t) = \frac{[\tau_{ij}(t)]^\alpha [n_{ij}]^\beta}{\sum_{k=allowed_k} [\tau_{ij}(t)]^\alpha [n_{ij}]^\beta} \quad (3.4)$$

A high value of α means that trail is very important and therefore ants tend to choose those paths as it means that other ants have also chosen it. On the other hand, low values of α make the algorithm very similar to a stochastic multi-greedy algorithm [8].

Ants choose the next city or node to travel to with the help of state transition rule. This rule is used to balance between **exploration** and **exploitation**.
if $q \leq q_0$ (exploitation)

$$S = arg_{u \in J_{(T)}} max[\tau(r, u)][\eta(r, u)]^\beta \quad (3.5)$$

Here q_0 is a constant parameter, q is a random variable, and S is the outcome of the probabilistic transition function.

Local updating rule: When ants move from one node to another in a single cycle, the pheromone trail between two nodes keep updating by different ants. This update of the pheromone trail is called local updating rule. The following formula denotes the rule:

$$\tau(r, s) \leftarrow (1 - \rho)\tau(r, s) + \rho\Delta\tau(r, s) \quad (3.6)$$

Here τ is a predetermined constant or function. The edge (r,s) is updated following each iteration of an ant search.

ρ is the evaporation constant. Here $0 \leq \rho \leq 1$

Global updating rule: In ACS, only the ants that successfully constructed the shortest paths let alone tour, from the very beginning of the process, are called the globally best ants and are allowed to deposit pheromone. This update is performed after all ants have completed their tours.

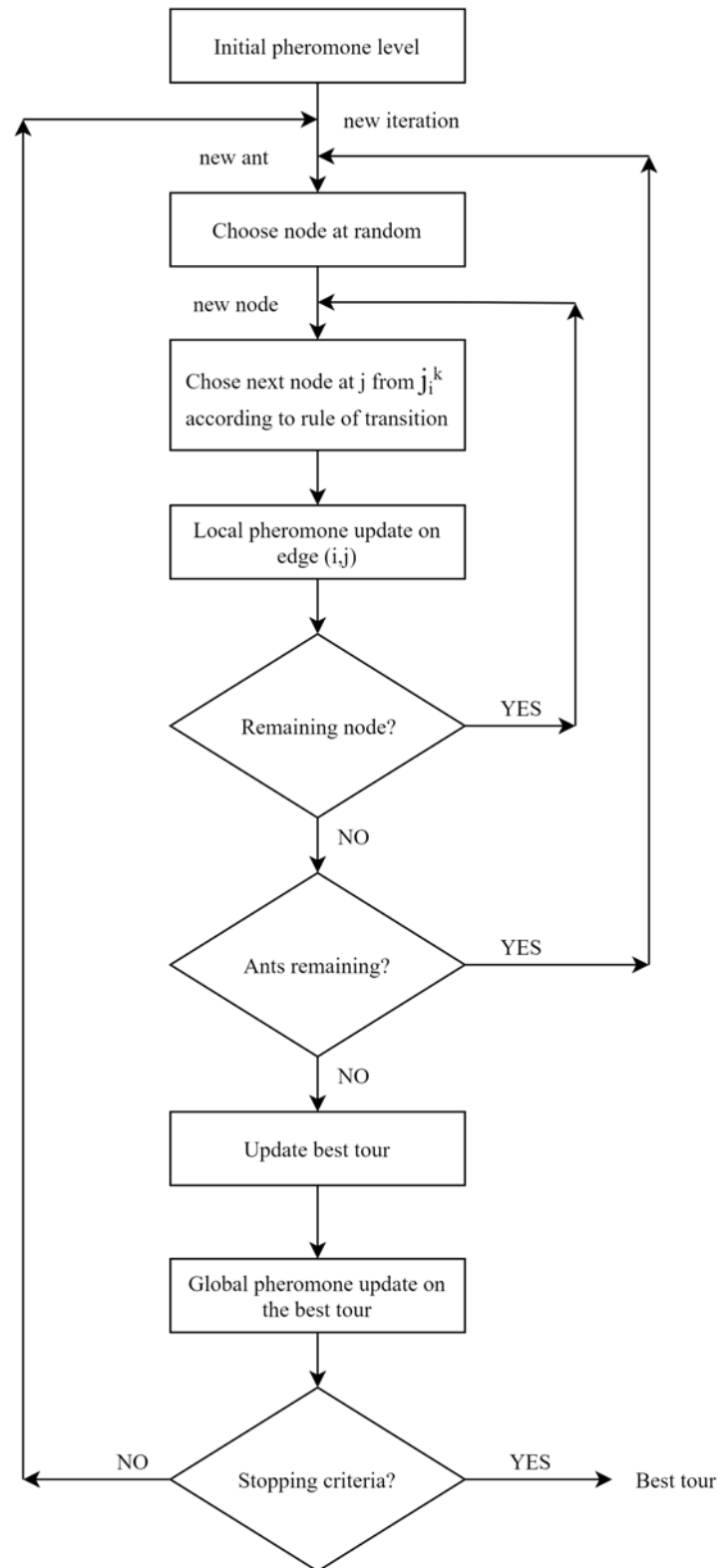


Figure 3.13: Flowchart of ACS Algorithm

Computational Complexity of the ACS:

The complexity of this ACS algorithm is $O(NC \times n^2 \times m)$ if the algorithm is stopped after $(N \times C)$ cycles, where number of cities and number of ants are denoted by n and m respectively.

Virtual Machine Consolidation using ACS Algorithm

The Service Level Agreements (SLAs) have formalized that the desired level of Quality of Service (QoS) has to be delivered to the customer by the server providers. The two most current used techniques are dynamic server provisioning and Virtual Machine (VM) consolidation. Dynamic server provisioning puts the servers in a switched off or a low power consuming mode, when the demand for cloudlets is low. On the contrary, When the demand of cloudlets is high the servers are put back on or put in high power consuming mode. On the other hand, Dynamic VM consolidation uses hardware virtualization technology. It shares a Physical Machine (PM) among multiple performance-isolated platforms called VMs and the sharing of PM resources is handled by the Virtual Machine Monitor (VMM). It leads to lower energy consumption. But to maximize resource utilization, the PM resource should be utilized properly. Ant Colony System (ACS) algorithm helps to overcome the problem. The simulation result of ACS-based VM consolidation (ACS-VMC) shows that it maintains the desired QoS while it reduces energy consumption. It also outperforms other VM consolidation approaches.

The Best Fit Decreasing (BFD) algorithm is used to allocate the initial VMs to PMs. This algorithm helps an efficient start. But the resource utilization of VMs continues to vary due to dynamic workloads. To adapt and optimize this workload, the ACS-VMC algorithm is proposed. This model consists of two agents. Local Agent (LA) and Global Agent (GA). LA observes the current resource utilization of the PM to solve the PM detection sub-problem. GA supervise and optimize the VM Placement by using ACS-VMC algorithm. As fig.(3.11) shows, LA monitors the CPU utilization and sets the PM into four different categories as normal, overload, predicted overload and under-loaded. Then the GA collects the data of every PMs from the LAs. By using the ACS-VMC algorithm it builds a global best migration plan. GA tells VMMs to perform the consolidation and VMMs perform the actual migration of VMs.

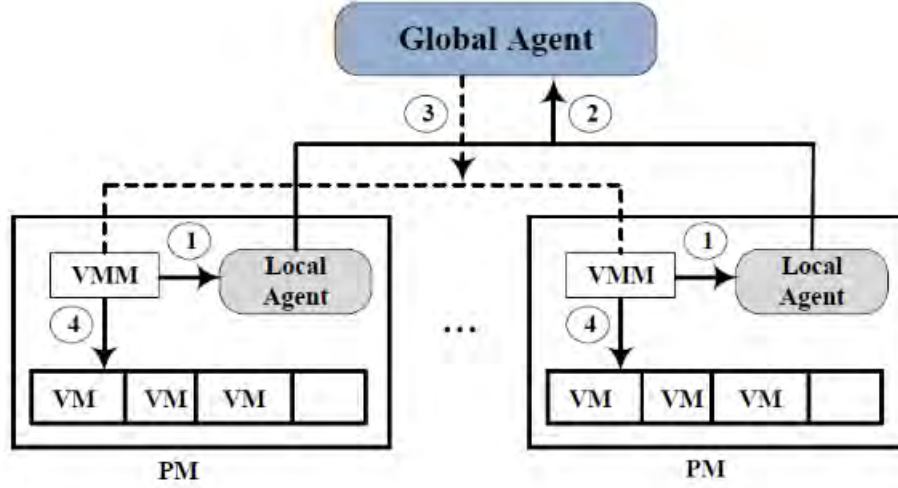


Figure 3.14: The system architecture of OEMACS

3.3.2 OEMACS Algorithm

We are designing an ACS with OEM operations to decrease the total amount of actively working servers and reduce power usage in cloud infrastructure. The resulting OEMACS algorithm looks for a result that has less numbers of host servers to accommodate all virtual machines. We identify the minimal number of servers as M_{min} in the global best solution S_{gb} . In the very beginning $M_{min} = N$ where N denotes the total number of virtual machines. It implies that S_{gb} 's original possible approach should be to install 'N' virtual machines in 'N' servers where one VM will correspond to one server.

Solution Construction:

In every single iteration, the algorithm OEMACS attempts to search and find an acceptable alternative in which there will be one host server less than M_{min} . Consequently, the agent ant attempts to put the 'N' virtual machines on the $M_t = M_{min} - 1$ servers.

The example of fig.(3.11) shows that the value of 'N' is 5 and ' M_t ' is 4 which are available. A $(M_t \times N)$ matrix is set as design scheme of the vertices. Here, every single vertex x_{ij} indicates a server assigned VM. In the corresponding columns, an ant's possible path is suggested by the non-directed arc in between two vertices.

It is to be noted that, every agent ant requires same processes in order to create a solution according to the design scheme shown in fig.(3.11). If we notice carefully, we can see that the virtual machines are shuffled arbitrarily prior to each cycle. After that the agent ant assigns VMs one by one to the servers to build a solution. We define the solution building process depending on one ant.

$$I_m = i | \sum_{n=1}^N x_{in} \times c_n + c_j \leq c_i \text{ and } \sum_{n=1}^N x_{in} \cdot m_n + m_j \leq M_i, 1 \leq i \leq M_t \quad (3.7)$$

$$\eta(i, j) = \frac{1.0 - \left| \frac{C_i - UC_i - c_j}{C_j} - \frac{M_i - UM_i - m_j}{M_i} \right|}{\left| \frac{C_i - UC_i - c_j}{C_i} \right| + \left| \frac{M_i - UM_i - m_j}{M_i} \right| + 1.0} \quad (3.8)$$

here UC_i and UM_i represent the usage of CPU and memory. OEMACS holds two factors designing the heuristic information:

1. the resource utilization for both CPU and memory
2. the balance of available resources

Increasing the usage of resources and utilizing multiple services in a balanced manner, facilitates consolidation of virtual machines.

$$\rho(i, j) = \frac{T(i, j)\eta(i, j)^\beta}{\sum_{k \in I_m} T(k, j)\eta(k, j)^\beta}, \forall i \in I_m \quad (3.9)$$

Here, $\beta > 0$ and it denotes the importance of heuristic information.

For VM_j , it selects server i from the server set I_m by applying the given state transition rule:

$$i = \begin{cases} \arg \min_{1 \leq K \leq M_t} \text{over}(K) & \text{if } q \leq q_0 \\ Q & \text{otherwise} \end{cases} \quad (3.10)$$

Here, Q is an integer selected randomly in $[1, M_t]$ according to the probability distribution formula mentioned below. If $q \leq q_0$ then the ant selects the server according to the minimum amount of overload rate after joining VM_j . The overload rate $\text{over}(i)$ is calculated as

$$\text{Over}(i) = \frac{|C_i - UC_i - c_i|}{C_i} + \frac{|M_i - UM_i - m_j|}{M_i} \quad (3.11)$$

Else, VM_j will be assigned according to the distribution rule of probability $q(i, j)$ to server i . Here $(1 \leq i \leq M_t)$. The probability distribution rule:

$$q(i, j) = 1 - \frac{\text{over}(i)}{\sum_{k=1}^{M_t} \text{over}(k)} \quad (3.12)$$

Fitness Function:

In every iteration, after an ant has completed constructing a solution the best way to evaluate the solution is to find its fitness value. With the help of the following hierarchical models, we evaluate the obtained solution:

$$f_1(s) = \begin{cases} \sum_{i=1}^{M_t} y_i & \text{if } S \text{ satisfies the capacity constraint} \\ M_t + 1 & \text{otherwise} \end{cases} \quad (3.13)$$

$$f_2(s) = \sum_{i=1}^{M_t} \left[\left(\frac{|C_i - UC_i|}{C_i} + \frac{|M_i - UM_i|}{M_i} \right) \cdot y_i \right] \quad (3.14)$$

Here, the amount of provided host servers in the on-going iteration t is denoted by M_t ; the probability of server i getting used in the final solution S is denoted by y_i thus $y_i = [0, 1]$. In eq.(3.13) if the obtained solution is feasible then the function $f_1(S)$ shows the total amount of servers that are active and not greater than M_t .

Eq.(3.13) gives the value of servers that we need to consolidate the process and eq.(3.14) gives the fitness value of the solutions by calculating the resources used. As such, the solutions with fewer servers and higher utilization is chosen in every iteration.

Management of Pheromone:

Pheromone on a path records the historical preference information. In the Ant-Colony Optimization process, two different types of pheromone updating rules are applied in order to update the pheromone left on the paths travelled. After an ant constructs a solution in a single iteration, each virtual machine pair (r, s) of the server i gets updated. The rule that is used to update the machines is:

$$\tau(r, s) \leftarrow (1 - \rho)\tau(r, s) + \rho\Delta\tau(r, s) \quad (3.15)$$

Here, $\rho = [0, 1]$ is the parameter of decaying pheromone.

$$\Delta\tau_i = \frac{1}{f_1(s_b)} + \frac{1}{NC_i + NM_i + 1} \quad (3.16)$$

$$NC_i = \frac{RC_i}{C_i} \quad (3.17)$$

$$NM_i = \frac{RM_i}{M_i} \quad (3.18)$$

Where, $\epsilon(0 < \epsilon < 1)$ is the pheromone enhance parameter, $f_1(S_b)$ is the number of physical machines i.e. servers used in S_b , NC_i NM_i represents the amount of CPU and Memory of server i that has been normalized and left over. Furthermore, RC_i and RM_i denotes the amount of CPU and Memory of server i that has been left over while C_i and M_i represents the total amount of CPU and memory of i .

The more occupies the server is, the more pheromone of VM pairs increases.

The Order Exchange and Migration Ant Colony System Algorithm

The OEMACS algorithm is composed of two separate algorithms. They are ACS (Ant Colony System) and OEM (Order Exchange & Migration). The full combined algorithm is given below:

Step 1: Initialization. The parameter for ACS is set to τ_0 . The feasible globally best solution is set as S_{gb} for placing N VMs on N servers. Thus, the number of

minimum servers is set to $M_{min} = N$. Set iteration $t = 1$ and maximum iteration as T .

Step 2: Set $M_t = M_{min} - 1$.

Step 3: The function of fitness $f(S)$ is applied in order to examine the fitness value of the constructed solution.

Step 4: The best solution S_b of the current iteration is set after evaluating the fitness value of the constructed solution. If S_b is feasible, S_{gb} is updated as S_b and $M_{min} = f_1(S_b)$ is set. Otherwise OEM local search is performed on S_b . S_{gb} and M_{min} are updated respectively if local search succeeds.

Step 5: After S_{gb} and M_{min} are locally updated, global pheromone update is eventually done on S_b and S_{gb} .

Step 6: Check if t is less than or equal to T . If not equal, then increment t by 1 and go to Step 3. Otherwise terminate the algorithm.

3.3.3 Collaboration of OEMACS & Cloud Computing

With the collaboration of OEM & ACS Algorithms, cloud computing has become surprisingly plausible and faster. Migrating idle machines from one server to another has never been this efficient. The combination of these two algorithms has given researchers the hope of something more useful and convenient.

i) **Order Exchange Operation:** The main objective of this operation is to enhance the resource utilization and share the resources of a non-overloaded server. To do so, every single VM of an overloaded server is compared with every single VM of the non-overloaded server with the view to make both the servers non-overloaded and balanced. The operation works in following way:

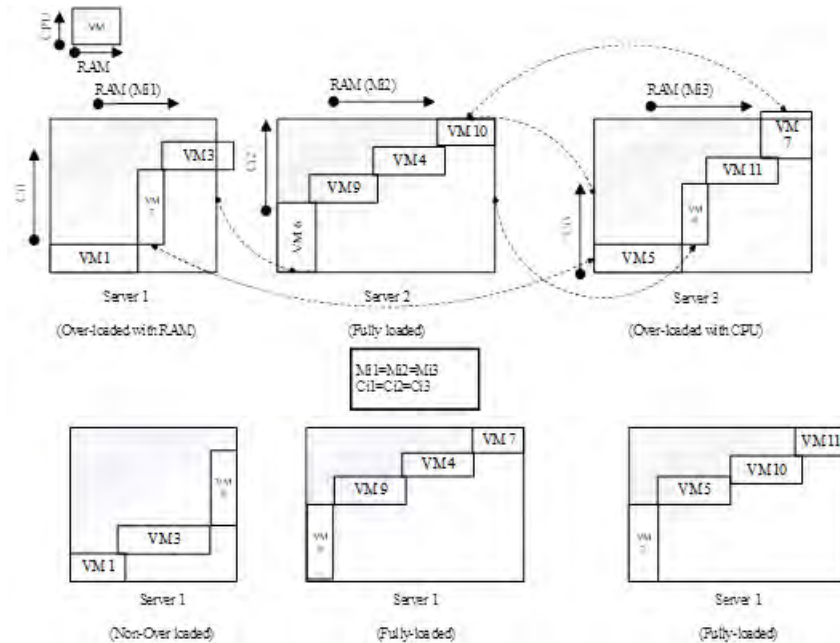


Figure 3.15: Order Exchange & Migration

After sorting is done in the overloaded server i , a server k which is non-overloaded

is selected and the Virtual Machines on that server are sorted out on the basis of the $|CPU - RAM|$ and in this case the contrast method is followed while giving preferences to the VMs. After that, every single VM in the queue of server i are compared with those of server k and the pair to exchange is determined.

The rectangles of different shapes denote the VMs inside servers. The dotted arrows show the possible exchange scheme of VMs to make the overloaded servers balanced and non-overloaded. In the example, 3 servers contain 11 VMs. The horizontal line of the VM rectangles represent the RAM size and the line drawn vertically represents the CPU size of all the Virtual Machines (VM). Server 1 and Server 3 are overloaded with RAM and CPU respectively. In order to make the two servers balanced and non-overloaded, VM 2,6,7,8 and 10 swapped their places with the other VMs in the non-overloaded servers.

ii) **Migration Operation:** This operation schedules the migration of a Virtual Machine from a server that is overloaded to a server that is non-overloaded and that can offer enough RAM and CPU to this VM. Before migration, the VMs in the overloaded server are checked one by one and compared with every single VMs on the non-overloaded server to determine which VM can be migrated to the destination server i.e the non-overloaded server. The operation of migration ends when all the overloaded servers becomes balanced or there is no significant VM movement possible. A possible migration scenario is illustrated in fig.(3.12), where VM 2,6,7,8, and 10 are migrated from their originating overloaded servers to destination non-overloaded servers.

3.4 Similar Algorithms

Power consumption by physical machines or servers is an important issue in cloud architecture. Latest studies have shown that servers in cloud data centers consume power that seems to be in linear statistical representation with the utilization of Central Processing Unit . An idle virtual machine can burn up to fifty to seventy percent of the electrical energy consumed by total number of virtual machines in a host server [2]. This scenario can be sketched using the following power model:

$$P(u) = k.P_{max} + (1 - k).P_{max}.u \quad (3.19)$$

Here, P_{max} is the highest amount of power consumed by a fully loaded host server
 k is the electrical energy consumed by an idle VM

u is the total Central Processing Unit utilization of the host

Although cloud is a virtual and non-physical network-based concept, its implementation includes physical and wired connectivity among servers, hard-drives and other network devices resulting in large number of data centers. As such power and energy consumptions increase total carbon footprint which is very hazardous for the environment. Green Cloud Computing is a concept of consuming less physical power and saving energy by making the infrastructure more virtual. With the help of the Ant Colony System (ACS) algorithm, it is possible to consolidate the Virtual Machines and make Cloud Computing more virtual than it already is. This algorithm describes a mechanism of indirect coordination, through the environ-

ment, between agents or actions e.g. communication about paths between ants with the help of pheromones. Using this technique, the Virtual Machines (VMs) that constitute Cloud Computing can be amalgamated and consolidated. Thus, Green Cloud Computing can be easily globalized.

The main objective of our research is to consolidate virtual machine and globalize green cloud computing. Virtualization is an extremely useful, efficient, and powerful technique that helps in the full utilization of hardware resources to its maximum capacity. With the help of this process, the resources of a physical machine or in this case, a server can be highly utilized by multiple number of virtual machines hosted in that server. In order to address the excessive power usage problem, VM consolidation is the one of the logical processes to apply. The power that is consumed by the virtual machines in a data center is nothing but electrical power and if we want to solve the high power consumption problem, all we have to do is to ensure that no unnecessary amount of electricity is being consumed by the machines. Not only that, along with power supply, we must make sure that the VMs are consuming the electricity in a way that the physical resources are utilized for serving application cloudlets and workloads.

Virtual Machine consolidation is a lengthy process, and it involves live migration, migration algorithms, favorable conditions, and environments and more. The whole process of virtual machine migration is dependent on some parameters like the destination host's capacity to store, host a new machine, the ability to assign resources to the VMs for completing cloudlets and more. At each stage of a server's life, it exhibits one or more environmental problems. So, in this paper, we are solely focusing on globalizing green cloud computing.

So, in order to make cloud computing safer, more efficient and environment friendlier, use of stigmergic algorithm is inevitable. With the use of an autocatalytic, multi-greedy heuristic algorithm, it is possible to keep the energy emission rate of servers bounded within a considerable range and thus we can ensure green computing.

After the very introduction of the Ant Colony System Algorithm in 1992, researchers and programmers all around the globe showed interest in solving the COP and NP-hard problems that could be solved using such algorithms only. They have used many similar EC-algorithms like Genetic Algorithm (GA), Particle Swarm Optimization (PSO), Multi-objective Ant Colony Optimization (MACO), Hybrid ACO and PSO and much more to address the VMP problem.

In all these algorithms, there remained some drawbacks and deficiencies which in return did not help in solving the problem in the most effective way possible. For such reasons, we have come up with a more efficient and rigorous EC-algorithm that will help address the VMP problem in a better way.

Chapter 4

Proposed Model

The model that we came up with is an algorithm to solve the Virtual Machine Placement problem. Our proposed model is called Active Idle Virtual Machine Migration Algorithm. With the help of this algorithm, we can easily migrate VMs from an overloaded server to a non-overloaded server and at the same time ensure the presence of maximum idle VMs in a destination server. With the help of this algorithm, the high-power consumption problem can be fixed and thus the carbon footprint of a server will come down to a great extent. Based on the objectives presented above, we designed a model that tends to optimize the utilization of resources used and energy consumed. These models try to minimize the resource wastage, SLA violations and excess energy consumption in the cloud infrastructure. The models are as follows:

$$\text{Min} \sum_{i=1}^n S_i \cdot SLA \quad (4.1)$$

$$\text{Min} \sum_{i=1}^n W_R = \sum_{i=1}^n [S_i \times \frac{|UC - UM|}{\sqrt{(UC)^2 + (UM)^2}}] \quad (4.2)$$

$$\text{Min} \sum_{i=1}^n P(u) = \sum_{i=1}^n [S_i \times (k \times P_{max} + (1 - k) \times P_{max} \times u)] \quad (4.3)$$

$$u = UC_i = \frac{\sum_{i=1}^n V_i \cdot UC_j}{\sum_{i=1}^n V_i}; 0 \leq V_i \leq 1; 0 \leq S_i \leq 1 \quad (4.4)$$

For S_i , possible constraints are

$$\sum_{i=1}^n S_i \cdot UC \leq R_{cpu} \quad (4.5)$$

$$\sum_{i=1}^n S_i \cdot UM \leq R_{Mem} \quad (4.6)$$

4.1 Active & Idle Virtual Machine Migration Algorithm

The algorithm with the help of which actively working virtual machines and idle VMs can be migrated from one overloaded or non-overloaded server to another non-overloaded server with the view to make power utilization efficient can be called Active Idle Virtual Machine Migration (AIVMM). Since our main problem regarding the Cloud Infrastructure and Virtual Machine Architecture is the massive energy consumption of the idle virtual machines which is around 50% - 70% of the total server input power, we plan to solve this problem by separating the idle VMs from the actively working VMs in a server. In order to do so, we need to design a model that can detect the number of actively working and idle VMs in a server at a given time. Virtual Machines act as independent, enclosed machines that use the resources they are allocated with during the process of virtualization. A VM is comprised of Processor (CPU), Memory (RAM), Space (Hard drive) and Network Bandwidth of its own. Based on the usage of CPU and RAM, we can easily detect which VM is actively busy and which VM is sitting idle. On using that information, it becomes quite simple to migrate all idle or active VMs from one server to another.

Our idea is to exchange all idle virtual machines of a single server with the actively working fully loaded ones of a non-overloaded server in order to ensure maximum number of active VMs in a single server at most of the times. As such, the actively working VMs will face no interruption in their power consumption since idle VMs tend to consume 50% - 70% of the total power of host server.

Judging from the CPU and RAM utilization of a VM, for an instantaneous time t , the following scenarios can co-exist in a cloud environment:

- $U_{m_j} > U_{c_j} [RAM > CPU]$
- $U_{m_j} < U_{c_j} [RAM < CPU]$
- $U_{m_j} \approx U_{c_j} [RAM \approx CPU]$

Here, U_{m_j} represents the memory utilization of a VM, U_{c_j} represents the processor or CPU utilization of a VM and $0 \leq U_{m_j} \leq 1, 0 \leq U_{c_j} \leq 1$ i.e U_{m_j} and U_{c_j} represents the percentage of RAM & CPU utilization.

4.2 Working Model

In the process of calculating total number of active and idle VMs of a single server at an instant t , we will use the resource utilization percentage of every single virtual machine. As a result, it will become easier to decide whether to migrate active VMs or idle ones. The possible scenarios of our applied method would be:

a) When $RAM > CPU$,

1. $U_{m_j} \gg U_{c_j}$
2. $U_{m_j} \approx U_{c_j}$

In (1), the RAM utilization is much greater than the CPU utilization which denotes that the VM is idle as it is using RAM or memory to keep itself on but not using much CPU compared to its RAM usage.

In (2), the RAM utilization is almost or exactly equal to the CPU utilization which denotes that the virtual machine is actively working with full load. Although the scenario of RAM usage being exactly equal to the CPU usage can be very rare.

b) When $RAM < CPU$,

1. $U_{m_j} \ll U_{c_j}$

2. $U_{m_j} \approx U_{c_j}$

When the CPU of a VM is working with full load but it is saving the data in cache or temporary memory rather than the RAM, then the CPU usage of the machine seems to be greater than the volatile memory i.e. RAM usage.

Here, in (1) the CPU utilization is much greater than the RAM utilization which denotes that the VM is active, but it is saving its data in the temporary memory or cache memory. A such the virtual machine is online and working with full load.

Furthermore, in (2) the RAM utilization is almost or exactly equal to the CPU utilization which denotes that the virtual machine is actively working with full load. Although the scenario of RAM usage being exactly equal to the CPU usage can be very rare.

Constraints of the functioning models (1), (2) and (3) are as follows:

For model (1), the possible constraints are:

- $|U_{m_j} - U_{c_j}| > 0$
- $0 \leq (U_{m_j} - U_{c_j}) \leq 1$
- $0 \leq (U_{m_j} - U_{c_j}) / (U_{m_j} + U_{c_j}) \leq 1, [U_{m_j} + U_{c_j} \neq 0]$

For model (2), the possible constraints are:

- $|U_{c_j} - U_{m_j}| > 0$
- $0 \leq |U_{c_j} - U_{m_j}| \leq 1$
- $0 \leq (U_{c_j} - U_{m_j}) / (U_{c_j} + U_{m_j}) \leq 1, [U_{m_j} + U_{c_j} \neq 0]$

For model (3), possible constraints are:

- $U_{m_j} - U_{c_j} = 0$
- $U_{m_j} - U_{c_j} = -(U_{c_j} - U_{m_j})$
- $|U_{m_j} - U_{c_j}| = |U_{c_j} - U_{m_j}|$

In order to calculate the total number of idle virtual machines in a server for an instant t , the following equations can be used:

$$V_0 = U_{m_j} - U_{c_j} \quad (4.7)$$

$$V_0' = \begin{cases} 0, & 0 \leq V_0 \leq 0.3 \\ 1, & 0.4 \leq V_0 \leq (0.9 \approx 1) \end{cases} \quad (4.8)$$

$$V_{idle} = \begin{cases} \sum_{i=1}^{N1} V_0', & [V_0' \in N1] \\ 0, & otherwise \end{cases} \quad (4.9)$$

Here, $N1$ is a set of virtual machines where $U_{m_j} > U_{c_j}$, V_0 represents the difference of RAM and CPU utilization while V_0' represents a piecewise function of V_0 . V_{idle} denotes the number of idle virtual machines in a server.

Again, in order to calculate the total number of working virtual machines in a server for an instant t , the following equations can be used:

$$V_a = U_{c_j} - U_{m_j} \quad (4.10)$$

$$V_a' = 1, if 0 \leq V_a \leq (0.9 \approx 1) \quad (4.11)$$

$$V_{active} = \begin{cases} \sum_{a=1}^{N2+N3} V_a' + c, & [V_a' \in (N2 + N3)] \\ c, & otherwise \end{cases} \quad (4.12)$$

$$c = \sum V_a \quad (4.13)$$

Here, $N2$ is a set of virtual machines where $U_{c_j} > U_{m_j}$, V_a represents the difference of CPU and RAM utilization of VM while V_a' represents a function of V_a . V_{active} denotes the number of actively working virtual machines in a server, c represents the number of active VMs from eq.(4.7)

When the number of actively working VMs and idle VMs at any time t , in any server i turns out to be equal, then the following equations can be used to calculate the total CPU and memory utilizations of the VMs:

$$UC_{V_a} = \sum_{j=1}^{V_{active}} U_{c_j} \times c_j \quad (4.14)$$

$$UM_{V_a} = \sum_{j=1}^{V_{active}} U_{m_j} \times m_j \quad (4.15)$$

Here, UC_{v_a} and UM_{v_a} denotes the CPU and Memory utilization of actively working virtual machines respectively, while U_{c_j} and U_{m_j} represents that of a VM. Also, c_j and m_j represents the total CPU and memory capacity of a VM.

$$UC_{V_i} = \sum_{j=1}^{V_{idle}} U_{c_j} \times c_j \quad (4.16)$$

$$UM_{V_i} = \sum_{j=1}^{V_{idle}} U_{m_j} \times m_j \quad (4.17)$$

The equations to calculate the total CPU and RAM usage of both an actively working and an idle virtual machine are respectively:

$$TU_{V_a} = UC_{V_a} + UM_{V_a} \quad (4.18)$$

$$TU_{V_i} = UC_{V_i} + UM_{V_i} \quad (4.19)$$

Here TU_{v_a} and TU_{v_i} represent the total utilization of both CPU and RAM by the active and idle VMs respectively, UC_{v_a} and UM_{v_a} represents the CPU and memory utilization of the active VMs while UC_{v_i} and UM_{v_i} represents those of the idle VMs.

The algorithm AIVMM is composed of ACS (Ant Colony System Algorithm) and OEM (Order and Exchange Migration) along with VMM (Virtual Machine Migration). The complete algorithm is sketched below:

Step 1: Initialization. Compute the number of active and idle VMs in a single host server. Take the number of active VMs as V_{active} and the number of idle VMs as V_{idle} . Compare V_{active} and V_{idle} . [Three possible scenarios can occur. Such as- i) $V_{idle} > V_{active}$, ii) $V_{idle} < V_{active}$ and iii) $V_{idle} = V_{active}$]

Step 2: The parameter for ACS is set to τ_0 . The feasible globally best solution is set as S_{gb} for placing N VMs on N servers. Thus, the number of minimum servers is set to $M_{min} = N$. Set iteration $t = 1$ and maximum iteration as T .

Step 3: Set $M_t = M_{min} - 1$. In each iteration, m ants construct m solutions and perform local pheromone updating on each solution.

Step 4: The fitness function $f(S)$ is applied in order to evaluate the fitness value of the constructed solution.

Step 5: The best solution S_b of the current iteration is set after evaluating the fitness value of the constructed solution. If S_b is feasible, S_{gb} is updated as S_b and $M_{min} = f1(S_b)$ is set. Otherwise OEM local search is performed on S_b . S_{gb} and M_{min} are updated respectively if local search succeeds.

Step 6: After S_{gb} and M_{min} are locally updated, global pheromone update is eventually done on S_b and S_{gb} .

Step 7: Check if t is less than or equal to T . If not equal, then set $t = t + 1$ and go to Step 3. Otherwise terminate the algorithm.

Step 8: After t terminates, calculate V'_{active} and V'_{idle} of the host server. If $V'_{active} = V_{active}$ and $V'_{idle} = V_{idle}$, then move forward to step 9. Otherwise update V'_{active} and V'_{idle} and then move to step 9.

Step 9: (a) If $V_{idle} > V_{active}$, then migrate all the actively working VMs from host server to a nearby (ACS) non-overloaded (OEM) server with the opposite scenario i.e. the server in which $V_{idle} < V_{active}$. The idle VMs of the destination server will be exchanged with the actively working ones from the host.

(b) If $V_{idle} < V_{active}$, then migrate all the idle VMs from host server to a nearby (ACS) non-overloaded (OEM) server with the opposite scenario i.e. the server in which $V_{idle} > V_{active}$. The actively working VMs of the destination server will be exchanged with the actively working ones from the host.

(c) If $V_{idle} = V_{active}$, calculate the CPU utilization ($UC_{v_a} \& UM_{v_a}$) and memory utilization ($UC_{v_i} \& UM_{v_i}$) of the VMs (both actively working and idle) and compare the total utilization of both type VMs using Eq.(4.18) & (4.19).

4.3 Salient features of AIVMM

4.3.1 A multi-objective approach towards VMM

Virtual Machine Migration (VMM) is a common phenomenon in cloud architecture. For the assurance of continuous connectivity and better performance, migration of virtual machine from one server to another is a must. In order to migrate a VM from one server to another and to accommodate the VM with required resource and enough power in the destination server, some parameters are needed to be satisfied. These parameters are:

- SLA (Service Legal Agreement)
- Resource Utilization
- Power or energy utilization

4.3.2 Service Level Agreement (SLA)

Maintaining the QoS (Quality of Service) is an important requirement in cloud computing. QoS requirements are commonly formalized as SLAs, which specify enterprise service-level requirements for data center in terms of minimum latency or maximum response time[28]. A workload independent metric called the SLA Violations (SLAV) is used to measure the SLA delivered in an IaaS cloud. It measures both violations due to Over-utilization (SLAV(O)) and violations due to Migrations (SLAV(M))[10]. Both these violations characterize the level of SLA Violations by the infrastructure independently and with equal importance.

$$SLA = \frac{\sum Req - AReq}{\sum Req} \quad (4.20)$$

Here, AReq denotes allocated requests

$$SLAV = SLAV(O) \times SLAV(M) \quad (4.21)$$

SLAV(O) represents the percentage of time during which active servers have experienced the resource utilization of 100%. It is defined as:

$$SLAV(O) = (M^{-1}) \sum_{i=1}^M \frac{T_{u_i}}{T_{a_i}} \quad (4.22)$$

Here, M is the number of servers, T_{u_i} represents total time server i has experienced the utilization of 100% leading to a SLAV; T_{a_i} is the total duration of server i being in the active state.

SLAV(M) represents the overall performance degradation by VMs due to migrations. It is computed as:

$$SLAV(M) = (N^{-1}) \sum_{j=1}^N \frac{C_{d_j}}{C_{r_j}} \quad (4.23)$$

Here, N is the total number of virtual machines. C_{d_j} is the performance degradation estimation of VM_j by migration, C_{r_j} is the total CPU capacity requested by VM_j in its lifetime.

4.3.3 Resource Utilization

To make full use of the resources, we have considered the CPU and the RAM to define the resource wastage in eq.(4.24) Here, W_R represents resource wastage, UC and UM denotes CPU and RAM utilization respectively. When W_R is large, more energy gets wasted.

$$W_R = \frac{|UC - UM|}{\sqrt{(UC)^2 + (UM)^2}} \quad (4.24)$$

4.3.4 Energy Utilization

Our focus is to minimize and utilize proper energy consumption as high use of energy creates high level radiation and increases pollution. To ensure green cloud computing our foremost target is to prevent idle virtual machines from disrupting the necessary power consumption for the actively working machines. Past researches have shown that idle virtual machines consume nearly 50% to 70% of the total server power and as such the neighboring virtual machines working with full load cannot get the salient power proportionately in time[11]. Thus, there remains a lag in service. The scenario can be depicted from the following model:

$$P = K \times P_{max} + (1 - k) \times P_{max} \times u \quad (4.25)$$

$$u = UC_i \tag{4.26}$$

Here, UC_i represents CPU Utilization of the server, P_{max} represents the peak power, k denotes the idle power consumption co-efficient.

4.4 Limitations of AIVMM

Although AIVMM algorithm helps to exchange active and idle virtual machines among servers and decrease unnecessary power consumption, this model is under development and in its testing phase. There are some limitations viable in the current model. They are:

- Time complexity of the algorithm is an important factor. The run-time complexity of AIVMM is quadratic and greater than that of OEMACS.
- Since this algorithm helps to migrate idle and actively working VMs from one server to other, it is a continuous process. An idle VM can start working anytime in the process, even when its being migrated while an actively working VM can cease to work and become idle any time in the process[29].
- In some cases, the differentiation between the resource utilization becomes quite complex. So, determining the number of active and idle VMs in a server can get complicated and problematic.
- Eq.(4.7) and (4.10) have some logical limitations. The equations are used to calculate the usage differences of RAM & CPU and CPU & RAM respectively. The equations do not clarify the scenario when $CPU = RAM \approx 0$.

Chapter 5

Result & Analysis

5.1 Result

We have successfully evaluated our algorithm- AIVMM in an advanced simulator CloudSim that helps to simulate cloud environments with necessary resources and infrastructure. The hardware device that we used in this purpose had the following configurations:

- Hardware - Intel Core i5-7200U 2.5 GHz Processor, 64-bit CPU, 1TB Hard Drive, 4GB DDR4 Memory
- Software - Eclipse IDE
- Environment - CloudSim 3.0.3 Language - Java
- OS - Windows 10 Education
- Network Bandwidth - 40 Mbps

We first formulized the AIVMM algorithm using the Eclipse IDE in Java. Then we collaborated the algorithm with OEMACS as AIVMM is a dependent algorithm. We conducted our experiments in CloudSim 3.0.3 that contain all the necessary classes required to simulate a cloud environment.

5.2 Analysis

5.2.1 Data Analysis of OEMACS & ACS Algorithm

Iteration	Max City Count	Max Ant Count	Highest Avg. Time (s)
150	4000	3200	2354.88
500	4000	3200	2641.52
1000	4000	3200	2932.66
2000	4000	3200	3905.34

Table 5.1: Dataset for ACS algorithm Simulation

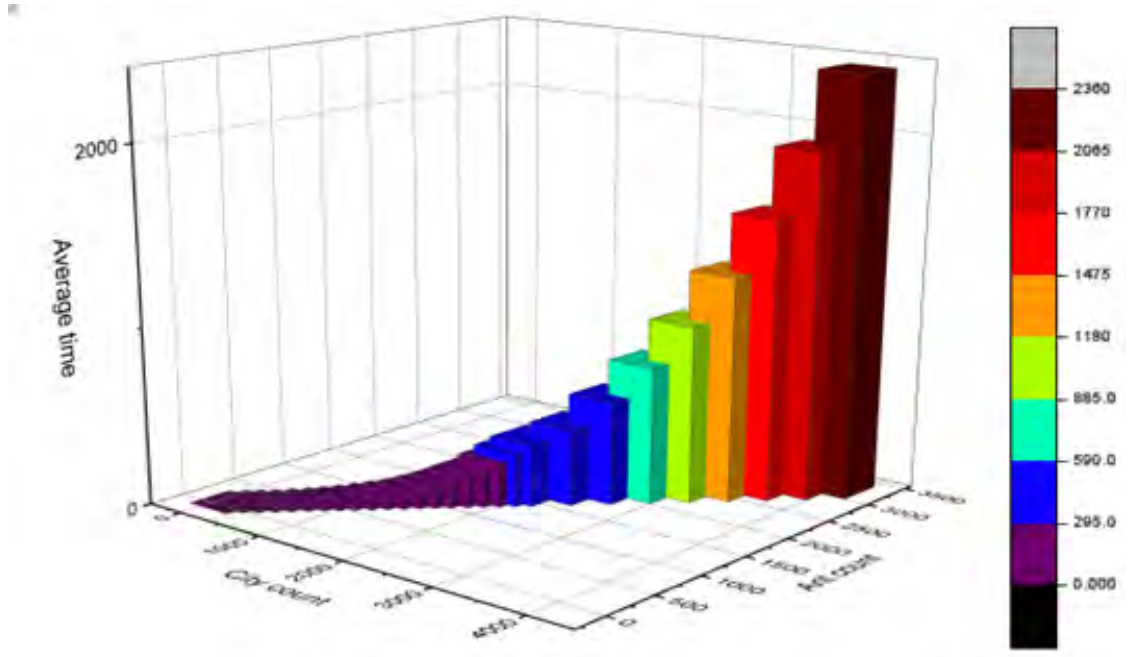


Figure 5.1: Average Time Complexity of OEMACS with Maximum Iteration 150

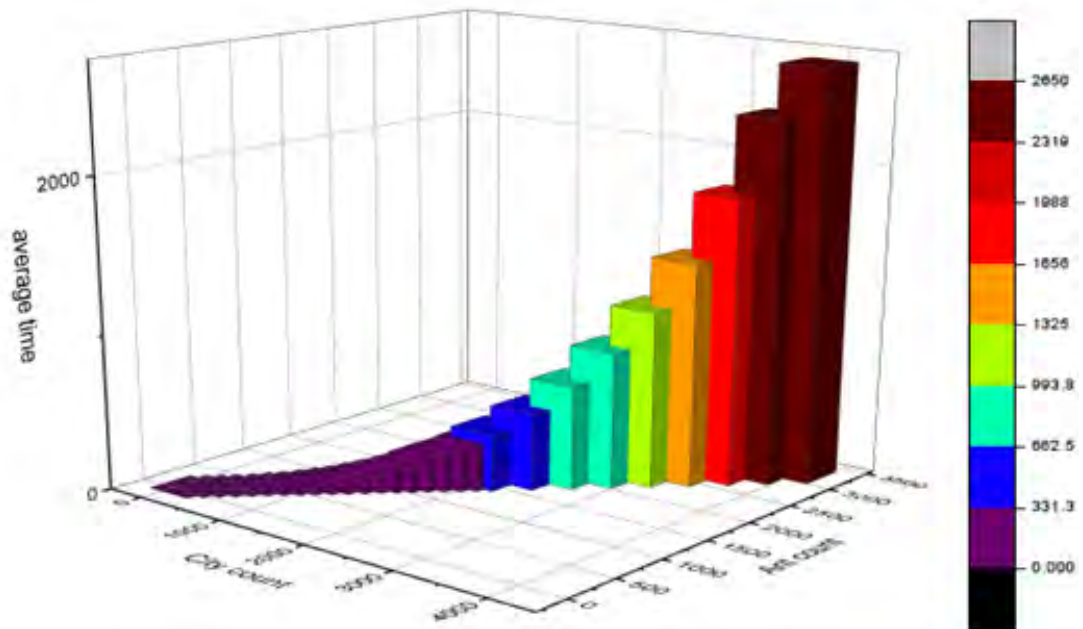


Figure 5.2: Average Time Complexity of OEMACS with Maximum Iteration 500

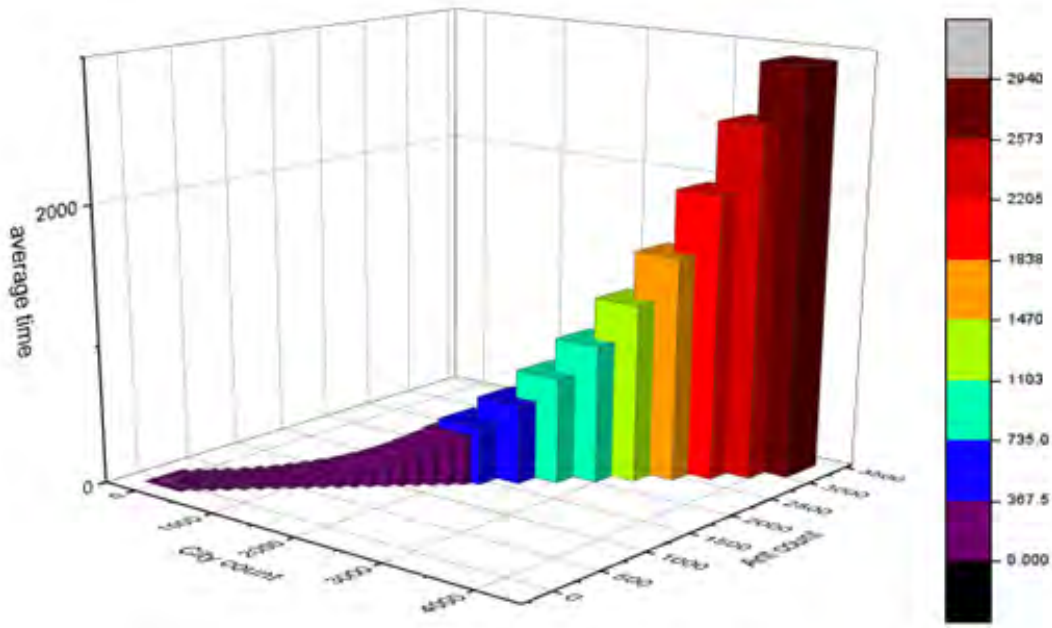


Figure 5.3: Average Time Complexity of OEMACS with Maximum Iteration 1000

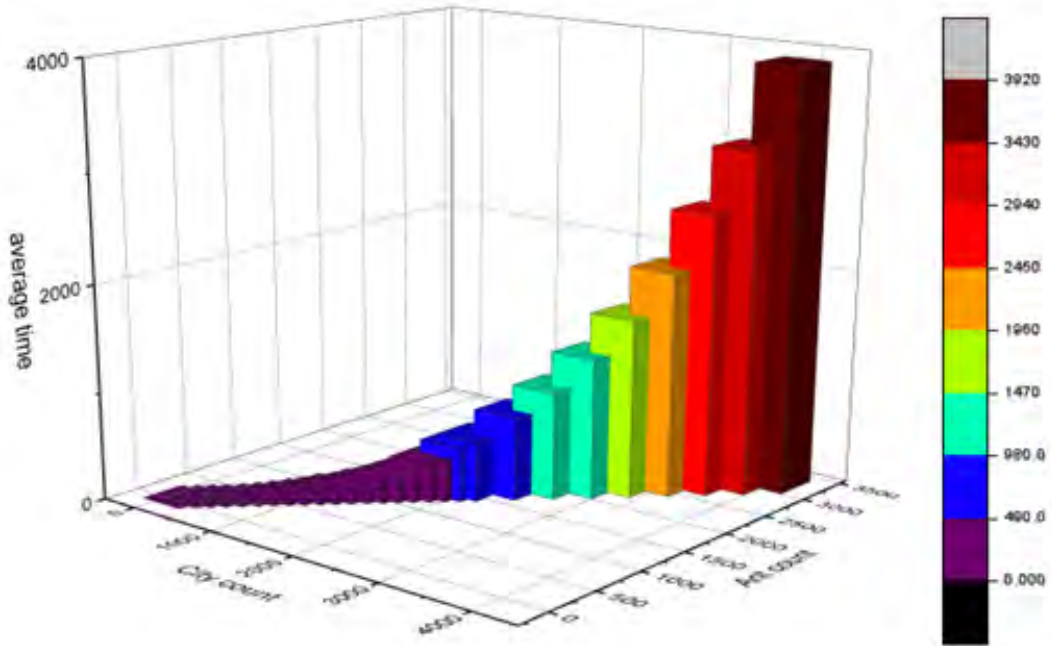


Figure 5.4: Average Time Complexity of OEMACS with Maximum Iteration 2000

5.2.2 Data Analysis of AIVMM Algorithm

We used the following datasets to represent the output of our simulation in which we tried to determine the relation between time and number of VMs added to the system after each time.

Time to instantiate (sec)	Number of VMs
1.1	1000
1.4	1200
1.6	1400
1.9	1600
2.2	1800
2.5	2000
2.8	2200
3	2400
3.4	2600
3.9	2800
4.1	3000

Table 5.2: Time to Simulation Instantiation

Amount of memory (kb)	Number of VMs
1000	1000
1100	1200
1210	1400
1311	1600
1412	1800
1513	2000
1615	2200
1716	2400
1817	2600
1918	2800
2019	3000

Table 5.3: Memory to Simulation Instantiation

Fig.(5.5) and (5.6) represents the amount of time and memory required respectively for the simulation instantiation when the number of virtual machines increases in a host server. The amount of time required is exponential to the number of VMs while the amount of memory required has a linear relationship.

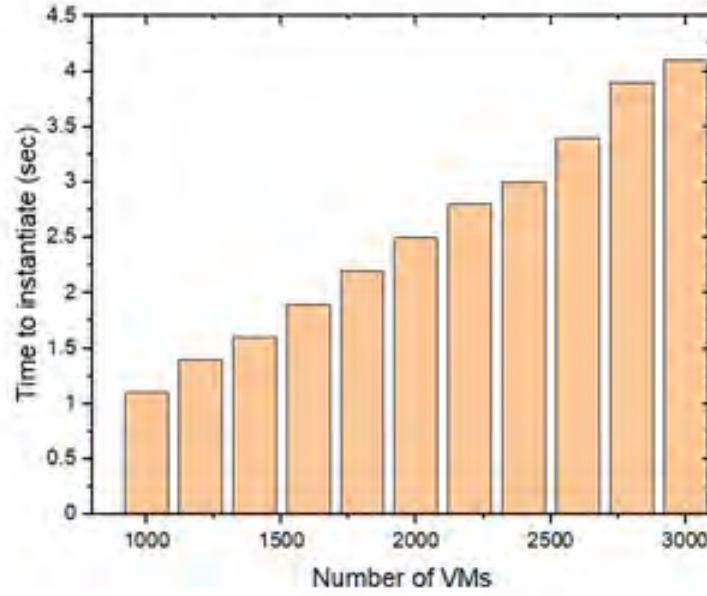


Figure 5.5: Time to Simulation Instantiation

In case of time overhead relating to the simulation instantiation, the growth with respect to time increases exponentially with the number of virtual machines. Here, the time required to insatntiate 3000 virtual machines is approximately 4 seconds, which is reasonable and standard for an experiment.

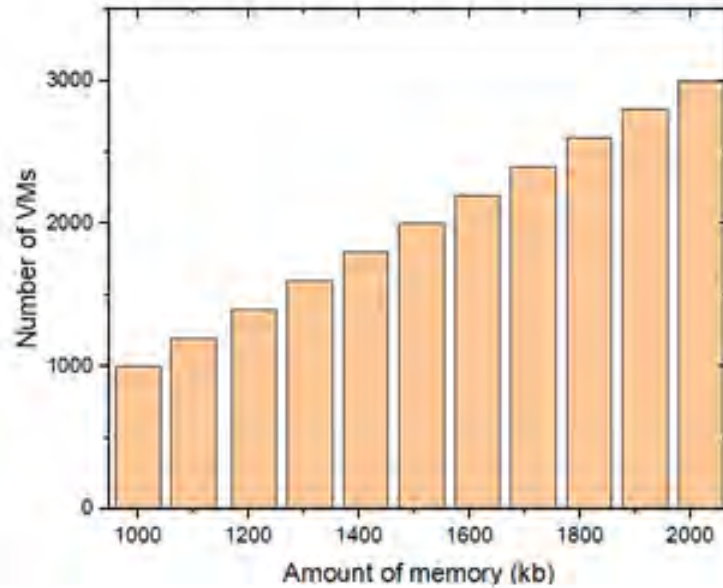


Figure 5.6: Memory to Simulation Instantiation

The growth in memory allocation is linear, with an experiment with 3000 virtual machines demanding approximately 2.1 MB RAM. It makes our simulation moderately suitable to run in desktop computers with standard processing powers.

We used the following dataset to represent the comparison between our parent algorithm- OEMACS and collaboration of OEMACS and AIVMM from the perspective of run-time.

VM ID	Time to migrate (sec) OEMACS	Time to migrate (sec) OEMACS+AIVMM
VM1	3.2	5.6
VM2	4.5	6.2
VM3	2.3	4.5
VM4	4.1	6.3
VM5	2.4	4.5
VM6	2.9	7.1
VM7	3.4	5.4
VM8	3.2	5.8
VM9	2.1	4.1
VM10	4.2	6.5
VM11	3.2	5.7
VM12	3.6	6.01
VM13	2.5	4.3
VM14	2.8	4.9
VM15	4.3	5.1

Table 5.4: OEMACS Migration

Fig.(5.7) denotes the comparison between our parent algorithm (OEMACS) and col-
laboration of all three algorithms mentioned in this paper (OEMACS and AIVMM).
The graph is sketched with time to migrate VMs against 15 virtual machines.

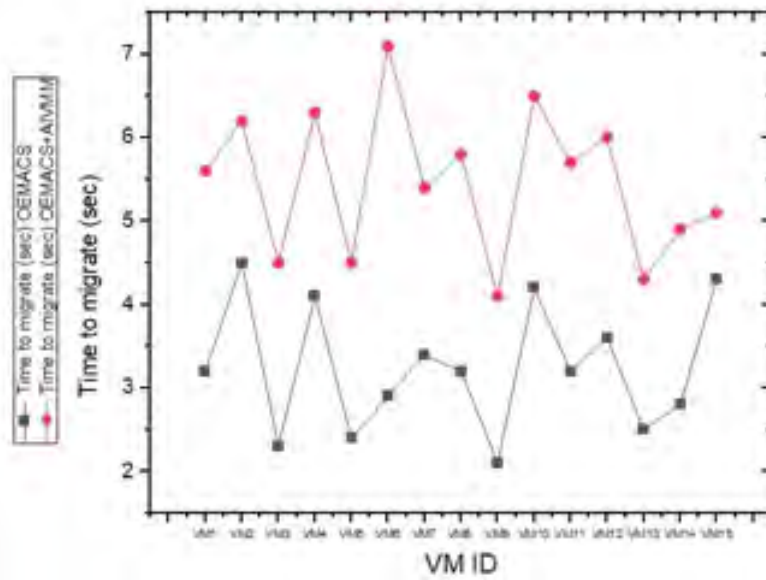


Figure 5.7: OEMACS Vs (OEMACS+AIVMM)

Chapter 6

Conclusion & Future Work

6.1 Conclusion

Virtual Machine Consolidation is a very salient feature of today's cloud infrastructure. With the help of continuous VM migration, cloud has become more resilient and lenient. But because of thermal throttle and excessive heat emission, the process of virtual machine migration has become a major issue. We have designed Active Idle Virtual Machine Migration algorithm which has not only solved the problem but has also consolidated the process to a great extent keeping the collateral parameters intact. Although there are limitations regarding run-time and migration threshold, in near future we plan to solve these issues by collaborating Migration of Idle Machine via Parking Server algorithm along with Active Idle Virtual Machine Migration algorithm.

6.2 Future Work

Our future plan for the AIVMM algorithm is to upgrade it by collaborating with another EC algorithm known as Migration of Idle Machines via Parking Servers (MIMPS) algorithm. In this algorithm, the idle virtual machines of a data center will be initially parked in a parking server, waiting to be assigned a task. After a VM gets assigned, it will switch its mode from idle and active and thus the system will migrate it back to its initial position. If the initial position does not remain vacant, the system will initialize OEMACS algorithm to migrate the VM to its nearest non-overloaded server. Thus, a balance will be created and the idle VMs will not consume unnecessary power or energy in their idle state.

Bibliography

- [1] M. Dorigo and L. Gambardella, “Ant colony system: a cooperative learning approach to the traveling salesman problem,” *Evolutionary Computation, IEEE Transactions on*, vol. 1, no. 1, pp. 53–66, 1997
- [2] A. Ashraf and I. Porres, “Using ant colony system to consolidate multiple web applications in a cloud environment,” in *22nd Euromicro International Conference on Parallel, Distributed and Network-based Processing (PDP)*, 2014, pp. 482–489.
- [3] M. Dorigo, V. Maniezzo, and A. Colnari, “The ant system: optimization by a colony of cooperating agents,” *IEEE Transactions on Systems, Man, and Cybernetics–Part B*, vol. 26, No. 2, pp. 29–41, 1996.
- [4] X. Liu, J. Zhang, “An Energy Efficient Ant Colony System for Virtual Machine Placement in Cloud Computing,” *Evolutionary Computation, IEEE Transactions on*, vol. 22, no. 1, pp. 30–65, 2018.
- [5] C. Yuan, and X. Sun, “Server Consolidation Based on Culture Multiple-Ant-Colony Algorithm in Cloud Computing,” *Sensors* 2019
- [6] L.M. Gambardella and M. Dorigo, “Solving symmetric and asymmetric TSPs by ant colonies,” *Proceedings of IEEE International Conference on Evolutionary Computation, IEEE-EC 96*, IEEE Press, 1996, pp. 622–627.
- [7] D. Fogel, “Applying evolutionary programming to selected traveling salesman problems,” *Cybernetics and Systems: An International Journal*, vol. 24, pp. 27–36, 1993.
- [8] R. Michel, M. Middendorf, An ACO algorithm for the shortest super sequence problem, in: D. Come, M. Dorigo, F. Glover (Eds.), *New Ideas in Optimization*, McGraw-Hill, London, UK, 1999, pp. 51–61.
- [9] D. Feng, Z. Wu, Z. Zhang, “A multi-objective migration algorithm as a resource consolidation strategy in cloud computing,” *Plos One* 2019
- [10] Y. Gao, H. Guan, L. Liu, “A multi-objective ant colony system algorithm for virtual machine placement in cloud computing,” *Journal of Computer and System Sciences* 79 (2013) 1230–1242
- [11] E. Feller, L. Rilling, and C. Morin, “Energy-aware ant colony based workload placement in clouds,” in *Proc. IEEE/ACM Int. Conf. Grid Comput.*, Lyon, France, 2011, pp. 26–33.

- [12] Y. Sahu, R. K. Pateriya, and R. K. Gupta, "Cloud server optimization with load balancing and green computing techniques using dynamic compare and balance algorithm," in Proc. IEEE 5th Comput. Intell. Commun. Netw. Conf., Mathura, India, 2013, pp. 527–531.
- [13] B. B. J. Suseela and V. Jeyakrishnan, "A multi-objective hybrid ACOPSO optimization algorithm for virtual machine placement in cloud computing," Int. J. Res. Eng. Technol., vol. 3, no. 4, pp. 474–476, 2014.
- [14] R. N. Calheiros, R. Ranjan and R. Buyya, "CloudSim: a toolkit for modeling and simulation of cloud computing environments and evaluation of resource provisioning algorithms," in Wiley Online Library (wileyonlinelibrary.com), 24th August 2010
- [15] Buyya R, Ranjan R, Calheiros RN. InterCloud: Utility-oriented federation of cloud computing environments for scaling of application services. Proceedings of the 10th International Conference on Algorithms and Architectures for Parallel Processing (ICA3PP 2010), Busan, South Korea. Springer: Germany, 21–23 May 2010; 328–336.
- [16] Buyya R, Ranjan R, Calheiros RN. Modeling and simulation of scalable Cloud computing environments and the CloudSim toolkit: Challenges and opportunities. Proceedings of the Conference on High Performance Computing and Simulation (HPCS 2009), Leipzig, Germany. IEEE Press: New York, U.S.A., 21–24 June 2009; 1–11.
- [17] Y. Tohidirad, S. Abdezadeh, Z. S. Aliabadi, "Virtual Machine Scheduling in Cloud Computing Environment", in International Journal of Managing Public Sector Information and Communication Technologies (IJMP ICT), Vol. 6, No. 4, December 2015
- [18] R. M. Chawda, O. Kale, "Virtual Machine Migration Techniques in Cloud Environment: A Survey", in International Journal for Scientific Research Development, vol. 1, issue 8, 2013
- [19] R. Ranjan and R. Buyya. Decentralized Overlay for Federation of Enterprise Clouds. Handbook of Research on Scalable Computing Technologies, K. Li et. al. (ed), IGI Global, USA, 2009 (in press).
- [20] J. Kommeri, T. Niemi, J. K. Nurminen, "Energy efficiency of dynamic management of virtual cluster with heterogeneous hardware", in International Journal for Scientific Research Development, vol. 5, issue 9, 2015
- [21] J. Kommeri, T. Niemi, O. Helin, "Energy efficiency of server virtualization", in International Journal On Advances in Intelligent Systems, vol. 5 (2012)
- [22] M. Dorigo, M. Birattari and T. Stutzle, "Ant Colony Optimization Artificial Ants as a Computational Intelligence Technique", Iridia – Technical Report Series: TR/IRIDIA/2006-023.
- [23] F. Salfner, P. Troger, A. Polze, "Downtime Analysis of Virtual Machine Live Migration", in The Fourth International Conference on Dependability, vol. 12, 2011

- [24] C. P. Sapuntzakis, R. Chandra, B. Pfaff, J. Chow, M. S. Lam, and M. Rosenblum, “Optimizing the migration of virtual computers,” *SIGOPS Oper. Syst. Rev.*, vol. 36, no. SI, pp. 377–390, 2002.
- [25] S. Akoush, R. Sohan, A. Rice, A. W. Moore, and A. Hopper, “Predicting the Performance of Virtual Machine Migration,” *Modeling, Analysis, and Simulation of Computer Systems, International Symposium on*, vol. 0, pp. 37–46, 2010.
- [26] I. K. Kim, S. Zeng, M. Humphrey, “A Supervised Learning Model for Identifying Inactive VMs in Private Cloud Data Centers”, in *Middleware Industry ’16*, December 12-16, 2016, Trento, Italy
- [27] S. N. Wang, H. X. Gu, and G. Wu, “A new approach to multi-objective virtual machine placement in virtualized data center,” in *Proc. IEEE 8th Int. Conf. Netw. Archit. Stor.*, Xi’an, China, 2013, pp. 331–335.
- [28] P. Patel, A. Ranabahu, and A. Sheth, “Service level agreement in cloud computing,” in *Proceedings of the Cloud Workshops at OOPSLA*, pp. 1–10, 2009.
- [29] H. Liu, C.-Z. Xu, H. Jin, J. Gong, and X. Liao, “Performance and energy modeling for live migration of virtual machines,” in *Proceedings of the 20th ACM International Symposium on High- Performance Parallel and Distributed Computing (HPDC ’11)*, pp. 171–181, June 2011.
- [30] F. Farahnakian, T. Pahikkala, P. Liljeberg, and J. Plosila, “Energy aware consolidation algorithm based on K-nearest neighbor regression for cloud data centers,” in *Utility and Cloud Computing (UCC), 2013 IEEE/ACM 6th International Conference on*, Dec 2013, pp. 256–259.