

Optimizing Workflow Task Allocation in Cloud Computing Environments

¹ Sreedevi Gangolu, ²Dr. Sumit Bhattachar

¹Research Scholar, Department of Computer Science, Singhania University, Pachari Bari, Jhunjhunu, Rajasthan.

²Associate Professor, Department of Computer Science, Singhania University, Pachari Bari, Jhunjhunu, Rajasthan

Abstract : Services like newline Infrastructure, Platform, and Software are made possible by cloud computing, an essential technology. It is difficult to distribute available resources and balance the load effectively while providing these services to newline's many cloud customers. There is no need to allocate additional resources to handle requests when using the defragmentation newline approach, which is one way for load balancing. This leads to greater resource utilisation. Lastly, an experimental investigation is conducted using the Cloud Sim tool to demonstrate the approach's applicability. The findings show that suggested mechanisms, such as task consolidation and defragmentation, outperform previous solutions.

Keywords: Infrastructure, Platform, Software, cloud.

INTRODUCTION

One of the newer methods of computing, "cloud computing" provides integrated computation, storage, and networking to end users via a subscription model. Similar to a subscription model, it offers consumers a great deal of flexibility in selecting certain services as needed.

The features offered by cloud computing have made it a popular choice among scientists. Software as a service (SaaS) involves the online deployment of software; platform as a service (PaaS) involves the provision of a computing platform that facilitates the easy and rapid design of web applications; and infrastructure as a service (IaaS) involves the distribution, on demand, of cloud computing infrastructure, including servers, storage, networks, and operating systems, rather than their purchase. While more and more organisations move their workloads to the cloud, and while cloud technology evolves at a dizzying rate to meet these demands, a major limitation of cloud computing is the scarcity of resources. To achieve the performance goals of infrastructure providers, cloud resource users, and applications, resource management involves distributing computing resources like virtual machines (VMs), storage, networking, and indirectly energy resources to a collection of applications. Management of resources encompasses a wide range of concerns. The scheduling of tasks is one of the primary concerns.

The primary concern in a cloud setting is task scheduling. Allocating application activities to appropriate resources while taking their interdependencies and independence into account helps meet QOS parameters, improves load balancing, maximizes resource utilisation, and reduces make-span. The tasks may be categorised as dependent or

independent depending on the degree of task dependence. Independent tasks are those that can be completed without interacting with any other activities. In the scheduling process, known as workflow, dependent tasks are distinguished from independent activities by virtue of the priority order in which they must be executed.

The workflow takes the directed acyclic graph (DAG) that represents the tasks of the apps into account. In a directed acyclic graph (DAG), each node stands for a task related to the issue, and each edge represents a dependency between tasks. Interactions between tasks in the process are possible. A basic workflow depicts actual work and includes a set of tasks, a sequence of activities, and the methods utilized to complete those tasks, whether they are performed alone or in a group. An example of a scientific workflow would be an application that relies on other activities, the execution of which might be rather complicated. Scientific workflows are the subject of this paper's study. In order to measure how well task scheduling algorithms operate, popular scientific procedures like MONTAGE, CYBERSHAKE, SIPHT, LIGO, and EPIGENOMICS are used as benchmarks. To test and refine the EM-HEFT method, the authors of this study used the LIGO and EPIGENOMICS procedures as reference cases.

LITERATURE REVIEW

Hatem Aziza et al (2020) Our world is becoming more and more dematerialized every day. The cloud is the source of all future technological innovations, the spark that ignites revolutionary concepts like blockchain and artificial intelligence. Datacenters, the backbone of the cloud, are responsible for over 20% of the world's total energy footprint. The host computers that make up these datacenters are

constantly improving in processing speed and power. In fact, even while doing nothing, the rising use of electrical energy is a direct result of the rise in computational capacity. Within this framework, we are discussing a development that is very taxing on the power grid and, by extension, the environment, as the internet is 1.5 times more polluting than air travel. It is very challenging to power using renewable energy because of the exponential consumption of digital. Since there are currently very few methods to restrict energy consumption, this challenge compels us to concentrate on solutions that optimise the utilisation of resources in the cloud environment to plan workflows. Reduced energy usage by cloud datacenters is an issue that this article helps to address. Dynamic Voltage and Frequency Scaling (DVFS) is a technology that may be used to limit the CPU frequency and impact energy usage; consequently, we are investigating it as a potential option. The power of the DVFS-based solution to optimise energy usage is shown by comparing it with other strategies. This is the main purpose of the paper.

Amit Agarwal et al (2014) The term "cloud computing" refers to a relatively new concept in distributed computing that allows users to pay for the resources they really utilize cloud is a network of interconnected computer systems that share resources like storage and processing power. The fundamental concept behind cloud computing is the effective provisioning of resources that are located in several physical locations. One of the numerous problems that the ever-evolving cloud is battling is scheduling. Scheduling is a collection of rules that govern the sequence in which a computer system executes its operations. An effective scheduler changes its approach to scheduling based on the job type and the environment. In this study, we compared our Generalised Priority algorithm to FCFS and Round Robin Scheduling, and we demonstrated how it efficiently executes tasks. If you want to see how well the method performs in comparison to more conventional scheduling algorithms, you should run it using the cloud Sim toolkit.

R. Singh, et al (2014) Cloud computing platforms rely heavily on task scheduling. Rather of basing task scheduling on a single criterion, users and cloud service providers must adhere to a complex set of rules and restrictions. The only thing this agreement is about is the level of service the customer expects from the service providers. Despite having a lot of things going on their end, providers have the critical responsibility of providing consumers with high-quality services in accordance with the agreement. To meet the objectives of tasks, we need to determine the best way to map or allocate their subtasks to the available resources (processors/computer machines). This is known as the task scheduling issue. In this research, we compare and contrast

several algorithms to see which ones work best in a cloud environment and which ones are the most flexible. Then, we attempt to suggest a hybrid method that might improve upon the current platform. In order for cloud service providers to be able to provide higher quality services.

CLOUD DATA RESOURCE ALLOCATION MODEL BASED ON ENERGY EFFICIENT LOAD BALANCER TECHNIQUE

Every single client may enjoy high reliability, resistivity, and improved cloud storage choices with the suggested energy efficient load balancing approach of the dynamic defragmentation model. On a smaller scale, distributed systems often make use of the Min-Min and Max-Min algorithms. The Max-Min scheduling algorithm is used in meta-assignments when the number of tiny tasks is greater than the number of large jobs. The makespan of the system is somewhat dependent on the number of little activities running concurrently with large ones. A modified version of the max-min method is suggested as a solution to this constraint. Below, we will examine the automatic dynamic defragmentation methodology that is used for load balancing.

Experimental Result and Discussion

The widely-used Java platform is used in the development of the software model. For the purpose of modelling cloud defragmentation, all of the modules have been fine-tuned. In order to replicate the suggested approach, a servlet-based application is created. The standard login handle given by a cloud provider is the main means of user authentication in the cloud. The data size is evaluated whenever the user requests to save any data chunk. Afterwards, a defragmentation timetable is set in motion by the cluster of servers after they assess the data amount and determine whether the cloud storage space is severely fragmented with files and requires defragmentation. Tomcat serves as the server backend and is built using open-source programming environments such as Eclipse JEE. In order to test our approach, we are making use of the free and open-source Eclipse JEE, which allows us to create interactive web applications for businesses.

Performance Analysis

Several assessment measures, including processing speed, throughput, instruction volume, and data volume, are used to assess the efficacy of the suggested energy efficient load balancing approach in the cloud storage model for automatic defragmentation. Assume for the sake of argument that the suggested technique works best when applied to meta-tasks with four tasks (T1, T2, T3, and T4) and when the scheduling manager is faced with a challenge involving two resources (R1 and R2). Table.2 shows the volume rate of rule predictions and information from T1 to T4, whereas Table.1

shows the speed and bandwidth range of each assigned resource. The task execution time (Te) and energy consumption (E) for each resource may be determined by

consulting Tables.1 and 2. Both the estimated and actual execution times of the jobs are shown in Table 3.

Table.1. Performance of the Max-Min and Modified Max-Min Technique

Problemsample	Resources	Processing Speed (MIPS)		Throughput (MBPS)	
		Max-Min	Modified Max-Min	Max-Min	Modified Max-Min
P1	R1	29	50	87	100
	R2	73	100	300	500
P2	R1	128	150	279	300
	R2	283	300	130	150
P3	R1	272	300	282	300
	R2	17	30	130	150

Table.1 shows how well the current max-min algorithm and the suggested modified max-min method performed. Here, R1 and R2 are the resources, and the problem sets are p1, p2, and p3. Both the current Max-Min method and the suggested improved Max-Min approach are shown in Figure.1, which displays their processing speeds.

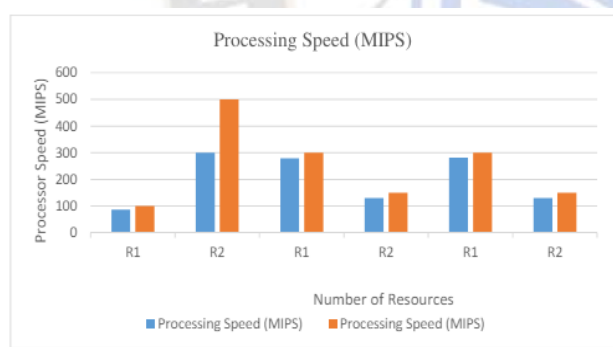


Figure.1 Processing Speed (MIPS)

Figure.1 shows that compared to the traditional Max-Min approach, the processing speed of the modified Max-Min

methodology is better. Processing speed is an area where the suggested modified Max-Min technique excels. Its speed is fast because it transmits a lot of data quickly.

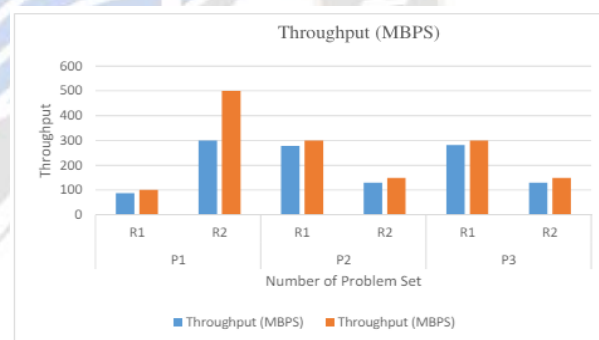


Figure.2 Throughput (MBPS)

Max-Min and Modified Max-Min throughput performance are shown in Figure 2. R1, R2 are the issue sets' resources, while P1, P2, and P3 are the problem sets' samples. Table 2 shows the data volume rates and instruction volume rates for both the current and proposed techniques.

Table.2. Performance of instruction and Data Volume

Problem sample	Task	Instruction Vol. (MI)		Data Vol. (MB)	
		Max-Min	Modified Max-Min	Max-Min	Modified Max-Min
P1	T1	112	128	39	44

	T2	52	69	53	62
	T3	201	218	83	94
	T4	11	21	43	59
	T1	212	256	73	88
P2	T2	21	35	25	31
	T3	298	327	89	96
	T4	201	210	432	590
	T1	11	20	63	88
P3	T2	321	350	27	31
	T3	192	207	83	100
	T4	17	21	37	50
	T1	11	20	63	88

You can see the amount of data and instructions needed for the job in Table.2. P1, P2, and P3 are the example problem sets here. There are four tasks (T1, T2, T3, and T4) in each issue set. The modified Max-Min approach outperforms the Max-Min in terms of performance. Nodes are exchanging data at a rapid pace. Time required for updated Max-Min to transmit massive amounts of data and instructions.

Table.3 Completion Time of Max min technique and Modified Max Min method.

Task/Resource	R1		R2		R3	
	Max-Min	Modified Max-Min	Max- Min	Modified Max-Min	Max-Min	Modified Max-Min
T1	152	102	250	236	178	123
T2	198	150	175	143	278	232
T3	98	58	210	188	256	193
T4	252	232	198	140	290	231

Task completion times for max-min and modified max-min with various resources are shown in Table 3. When compared to the conventional Max-Min method, the value of (T) from the Modified Max-Min approach indicates superior results.

Once the algorithm's anticipated time (E) has been determined, the submitted work on each feature is prioritized, and the job with the highest expected time (E) is given to the resource with the lowest total time (T).

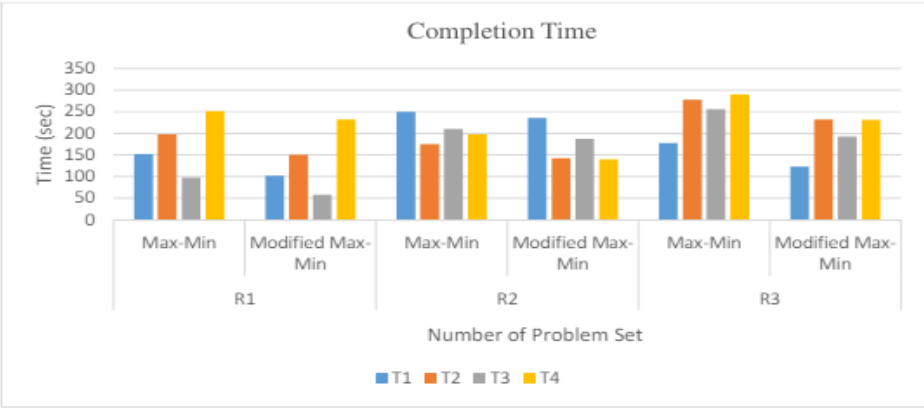


Figure.3 Performance of Completion Time

The new Max-Min method and the current method are shown in Figure.3 for the job completion time. The y-axis shows the total number of problem sets, while the x-axis shows the time in seconds. The sets of resources for the three problem sets

are R1, R2, and R3, correspondingly. Tasks are completed much more quickly and with less effort using the suggested modified Max-Min methodology as opposed to the conventional Max-Min method.

Table.4. Performance of Makespan

Problem Sample	Max-Min	Modified Max-Min
P1	11	10
P2	9	8
P3	5	4

You can see the makespan of each issue set in Table 4. The total duration of the rules is called the makespan. In a dynamic scheduling system, the problem is described as the

processing of each work and the use of online results to determine which jobs are assigned before assuming ultimate execution.

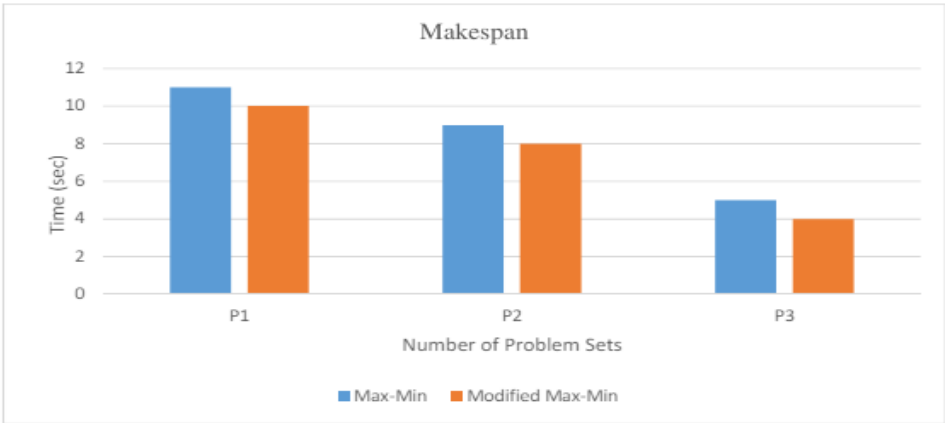


Figure.4 Performance of the Makespan

Both the traditional Max-Min and the modified Max-Min methods are shown in Figure.7, which displays their makespan performance. The results of comparing the original max-min method with the significant max-min algorithm, as

illustrated in Figure.7, are shown in Table.4.

MINIMIZING ENERGY CONSUMPTION IN CLOUD DATACENTERS AND VIRTUAL MACHINES USING TASK CONSOLIDATION

A service level agreement (SLA) between a cloud provider and its customers is the basis for cloud computing, which is a platform for large-scale distributed, parallel computing. The bulk of organisations are outsourcing their IT administrations to the cloud, which means the CSPs are using vast processing capacity. In addition to being more expensive, DCs discharge greenhouse gases like carbon dioxide into the atmosphere because to the large amounts of electrical energy they utilises. These problems are addressed by using an efficient assignment consolidation mechanism. Conversely, in order to minimise electric power consumption without breaching the SLA, dynamic VMC is used using the live migration approach. Turning off hosts in DC that are operating

uncertainly is also helpful.

EXPERIMENTAL RESULT AND DISCUSSION

This experimental effort makes use of the Cloud Sim tool, a Mac OS X device, 4 GB of RAM, and an i5 core processor 2.5GHz. The Net Beans IDE and its toolkit are designed for use with a single DC. While running the simulation, the virtual machine and host are set up similarly to the task consolidation model that was suggested. The table displays the results of the simulation comparing the energy efficacy of the suggested task consolidation technique with the RR method used in the Cloud Sim tool.

Table.5. Energy Consumption of RR and Task Consolidation

Cloudlet Number	Task Consolidation	RR
4	0.07	0.11
8	0.15	0.18
16	0.30	0.37
24	0.44	0.58
32	0.59	0.77

Table 5. shows how various numbers of cloudlets affect energy consumption performance. The energy usage goes up

as the number of cloudlets goes up. Energy consumption is lower when using the RR Task Consolidation method.

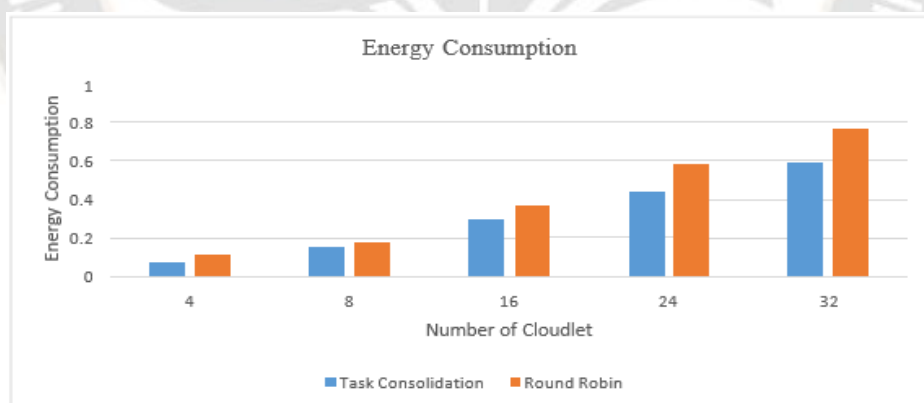


Figure.5. Energy Efficiency between RR and Task

Consolidation

Figure.5 is a graphical depiction of the performance metrics for energy consumption for both the current RR approach and the suggested task consolidation. The Y-axis shows the energy consumption, while the X-axis shows the quantity of

cloudlets. Compared to the RR, the task consolidation approach yields superior results. By combining an LRR method with a MMT policy as the virtual machine selection policy, it is possible to achieve much better energy usage with a significantly a smaller number of migrated virtual machines, outperforming existing dynamic VMC techniques.

Table.6. Energy Efficiency of existing LRR-MMT and Modified LRR-MMT methods

Number of Cloudlet	Modified LRR-MMT (proposed)			LRR-MMT (existing)		
	VM Migrated	Energy Consumed	Execution Time (ms)	VM Migrated	Energy Consumed	Execution Time (ms)
16	5	0.07	3.2	13	0.10	6.4
24	8	0.08	4.8	22	0.14	9.6
32	14	0.12	6.4	27	0.17	12.8

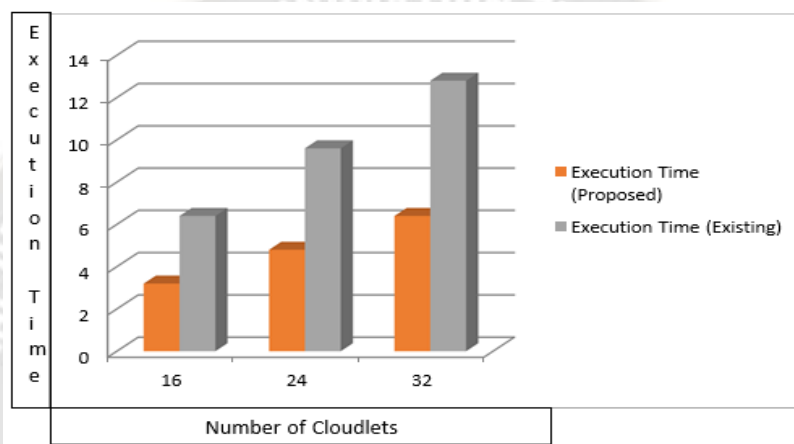


Figure.6 Execution Time between Proposed and Existing methods

The results of the improved LLR-MMT and LRR-MMT approaches to virtual machine migration and energy usage are shown in table 6. Table.6 demonstrates that the simulation results reveal that the modified LRR method outperforms the current LRR technique of Cloud Sim when the task consolidation approach is used to the LRR technique. By implementing changes to the task scheduler of the LRR-MMT technique in Cloud Sim using the task consolidation approach, we were able to reduce the number of virtual machines that needed to be migrated, leading to less energy consumption in the data centre. These results were compared to those of the original LRRMMT approach in Cloud Sim. All of the experimental analysis was done in just one DC. The data center's 20 hosts were home to 20 virtual machines. Each

server has a single-core central processing unit (CPU) with a random allocation of 1000, 2000, or 3000 MIPS of processing power, and 10,000 MB of random-access memory (RAM) and 100,000 MB of random-access bandwidth (bandwidth). Virtual machines typically use 250, 500, 750, and 1000 MIPS of CPU time, respectively, and have 2500 MB of bandwidth and 128 MB of RAM.

The findings show the power efficiency and the number of virtual machines (VMs) migrated using MMT, and the experimental investigation was conducted using the XEN hypervisor. The table and preceding sections detail the various VM selection policy techniques. Figure 7 shows the results of several virtual machine allocation policies.

Table.7. Performance of different VM Allocation Policy

VM Allocation Algorithm	Number of VMs Migrated	Power Consumed (KWH)
LR	13	0.07
IQR	19	0.11
MAD	13	0.09

LRR	15	0.12
DVFS	0	0.33

Energy usage and the number of VM migrations are shown in Table.7, which compares the various virtual machine allocation strategy techniques. Energy usage is 0.07 kwh for the LR method and 0.09 kwh for the MAD method, but both demonstrate an identical number of virtual machine

migrations. The AM IQR approach uses 0.11 kwh of energy, whereas the LRR method uses 0.12 kwh, and both reflect migrations of 19 and 15 virtual machines, respectively. Lastly, the DVFS technique uses 0.33 kwh of energy but does not include VM migration.

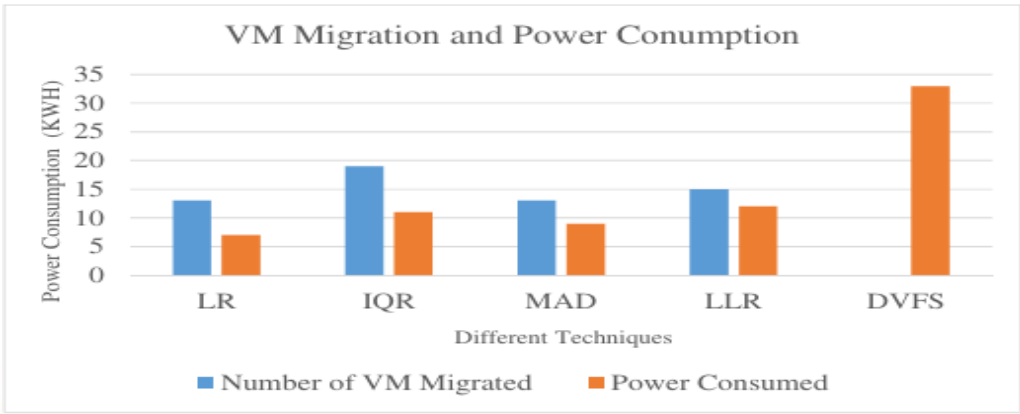


Figure.7. Performance of Power consumption and VM Migration

Different virtual machine allocation policies, including LR, IQR, MAD, LLR, and DVFS, are shown in Figure.7. All of the methods demonstrate a large number of VM migrations with little energy use. Despite having the greatest energy

usage, the DVFS approach exhibits zero VM migration. Based on experimental data, it seems that the LR + MMT combination produces better outcomes than alternative strategies.

Table.8. CPU Utilization Performance of different VM Allocation Technique

VM Migration	CPU Utilization (%)				
	LR	IQR	MAD	LLR	DVFS
5	92	90	89	97	90
10	83	85	78	83	80
15	77	73	72	77	70
20	64	68	61	64	60
25	53	57	50	53	48

Performance of CPU utilisation vs various VM migration and allocation rules is shown in table 7. With regard to the migration of 10 virtual machines, the LR and LLR account

for 83% of the CPU utilisation. Out of all the policies, the DVFS policy uses the CPU the most.

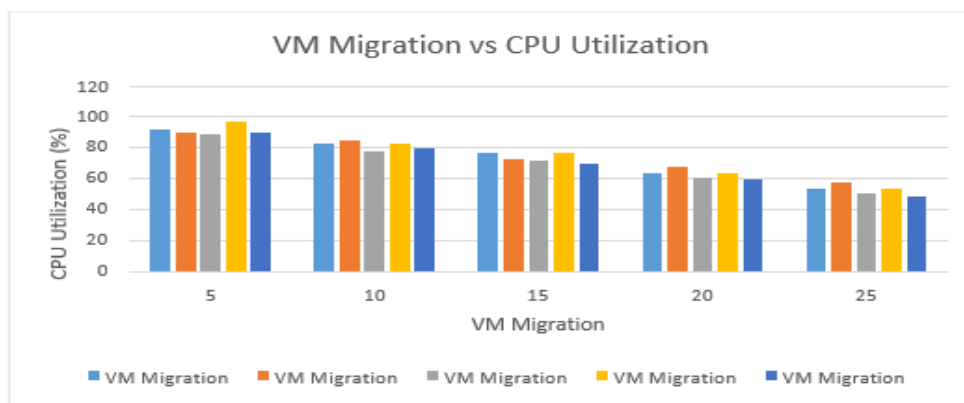


Figure.8. Performance of CPU Utilization

Different virtual machine allocation rules are graphically shown in Figure.8 to show how they perform in terms of CPU utilisation. Each virtual machine allocation policy minimizes power consumption and maximizes CPU utilisation without breaching the service level agreement. The various virtual machine allocation strategies help to maximize resource utilisation and decrease waste in the data centre.

CONCLUSION

Cloud computing refers to a platform for large-scale distributed and parallel computing where consumers sign service level agreements (SLAs) with cloud service providers and in exchange get computer resources as a utility. The majority of companies are moving their IT administration to the cloud in order to lower the initial capital expenditure in IT framework setup and to reduce the burden of equipment and programming maintenance. By default, when you upload an app to the cloud, it should distribute the processing burden evenly.

REFERENCE:

- [1] Abazari F, Analoui M, Takabi H, Fu S. MOWS: multi-objective workflow scheduling in cloud computing based on heuristic algorithm. *Simulation Modelling Practice and Theory*. 2019 May 1;93:119-32.
- [2] Agarwal, Amit & Jain, Saloni. (2014). Efficient Optimal Algorithm of Task Scheduling in Cloud Computing Environment. *International Journal of Computer Trends and Technology*. 9. 10.14445/22312803/IJCTT-V9P163.
- [3] Alawad NA, Abed-alguni BH. Discrete island-based cuckoo search with highly disruptive polynomial mutation and opposition-based learning strategy for scheduling of workflow applications in cloud environments. *Arabian Journal for Science and Engineering*. 2021 Apr;46(4):3213- 33.
- [4] Aziza, Hatem & Krichen, Saoussen. (2020). Optimization of workflow scheduling in an energy-aware cloud environment. 1-5. 10.1109/OCTA49274.2020.9151653.
- [5] Bharathi S, Chervenak A, Deelman E, Mehta G, Su MH, Vahi K. Characterization of scientific workflows. In *2008 third workshop on workflows in support of large-scale science* 2008 Nov 17 (pp. 1-10). IEEE.
- [6] Chen WN, Zhang J. An ant colony optimization approach to a grid workflow scheduling problem with various QoS requirements. *IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews)*. 2008 Oct 31;39(1):29-43.
- [7] Ismayilov G, Topcuoglu HR. Neural network based multi-objective evolutionary algorithm for dynamic workflow scheduling in cloud computing. *Future Generation computer systems*. 2020 Jan 1;102:307-22.
- [8] Manasrah AM, Ba Ali H. Workflow scheduling using hybrid GA-PSO algorithm in cloud computing. *Wireless Communications and Mobile Computing*. 2018 Jan 1;2018.
- [9] Mangalampalli S, Pokkuluri KS, Kocherla R, Rapaka A, Kota NR. An Efficient Workflow Scheduling Algorithm in Cloud Computing Using Cuckoo Search and PSO Algorithms. In *Innovations in Computer Science and Engineering 2022* (pp. 137-145). Springer, Singapore.
- [10] Mangalampalli S, Pokkuluri KS, Satish GN, Kumar KV. An Effective Workflow Scheduling Algorithm in Cloud Computing Using Cat Swarm Optimization. *ECS Transactions*. 2022 Apr 24;107(1):2523.
- [11] Sadiku MN, Musa SM, Momoh OD. Cloud computing: opportunities and challenges. *IEEE potentials*. 2014 Jan 9;33(1):34-6.
- [12] Saeedi, Sahar, et al. "Improved many-objective particle swarm optimization algorithm for scientific workflow scheduling in cloud computing." *Computers & Industrial Engineering* 147 (2020): 106649.
- [13] Singh, Raja Manish et al. "Task Scheduling in Cloud Computing : Review." (2014).