

# Energy Efficient Task Scheduling in Cloud Data Center

Deepanshu Kumar<sup>1\*</sup> and Sudhanshu Kulshrestha<sup>2</sup>

<sup>1</sup>Jaypee Institute of Information Technology, Noida, Uttar Pradesh, India. Email: [dkumar247@gmail.com](mailto:dkumar247@gmail.com)

<sup>2</sup>Assistant Professor, Jaypee Institute of Information Technology, Noida, Uttar Pradesh, India.

Email: [sudhanshu.kulshrestha@jiit.ac.in](mailto:sudhanshu.kulshrestha@jiit.ac.in)

\*Corresponding Author

**Abstract:** Cloud computing is emerging as a necessary need for the IT industry in order to reduce the setup and operational cost of its infrastructure. There is a huge requirement of computing resources to satisfy customer demands. A minute delay in a service may result in a measurable amount of loss for an organization. Response time is a major metric for evaluating performance of cloud applications. Cloud data centers form backbone of cloud computing. Data centers consume enormous amount of energy. Server racks have processing units, storage and network interface. Energy is dissipated at the server racks and cooling units. Various task scheduling algorithms and virtual machine scheduling algorithms have been proposed to measure the loss in performance but the energy loss is kept at the lowest priority. The paper is focused on discussing about the two techniques that maintain a scheduled routine for tasks arriving in a data center through a simulation scenario. VM-specific scheduling of tasks is done for assignment of the tasks to single or multiple virtual machines. Comparison of the two techniques, time-shared and space-shared technique is also done to give the reader a clear view about the situation in which both techniques are used. Future work is also discussed in the same context.

**Keywords:** Cloud computing, Cloudlet, Energy, Space-shared, Time-shared.

## I. INTRODUCTION

Cloud computing as a whole is a vast model used for providing computing resources to people who do not have enough storage to keep the resources, or who cannot afford to keep the resources. The model mainly benefits by its great pay per use feature which gives the user flexibility to use a resource more and paying for the price according to the amount of usage just like an e-bill or a water bill. So, according to NIST [29], cloud computing is a model that enables on-demand network

access to a shared pool of configurable resources like networks, servers, storage, applications and services that can be rapidly provisioned with minimum effort given by the service provider. This also gives a view that what can be the resources provided to the user like, not only computing resources are equipped by the user, but even network equipment play an equal and essential part in making the communication reliable for the user. First layer being the Software as a Service (SaaS), where most of the processing takes place on a server provided by the service provider and to interact with the processed data, an application is needed. Google Apps and Salesforce are valid examples of SaaS. Network management and Web framework teams are needed separately for maintaining the stack. Often, applications built on this scale face difficulty in expanding according to user demands, also making it difficult to access around the world. Platform as a Service (PaaS) is the layer which provides a single platform to build cloud applications through which application is deployed faster, and even much less electric power is wasted instead of having multiple servers in same place. Users can access custom-built apps over the internet and focus more in innovative ideas to enhance the application. Examples can be IBM Bluemix or Microsoft Azure. Cloud is more dependent on hardware, more of virtual connections are needed for the users to interact with the same hardware components. Bandwidth and IP addresses are also provided as they are crucial. Hosting a website on virtual servers can be an example, where these virtual servers are a part of the physical servers kept in data centers. Infrastructure as a Service (IaaS) hold the service provider responsible for scalability and maintenance of hardware. Notice Fig. 1 for more clarification.

Servers perform a core task to store and transfer information in cloud computing. This covers a larger part of net power consumption. According to [9], the efficient management of these servers also result in minimizing operating cost by turning down servers that have lower load and transferring the load to ones having higher load, this term is also referred as dynamic rightsizing.

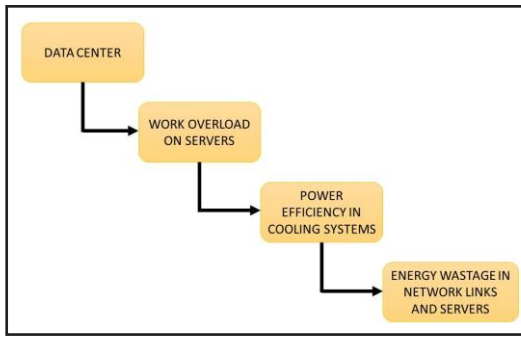


Fig. 1: Whole Picture Showing Problems

### A. Motivation

One of the main power absorbing component in cloud computing is the CPU. In coming time, user count will increase and there will be an obvious need to limit the usage of power at a particular level. Further, this will lead to a need for energy efficient servers. A server consists of high valued resources like RAM, secondary memory and network components. Even a small data center has the capability to save huge amount of money for the service providers by saving energy if it operates on efficient equipment. Critical issues like excessive heat collected in a data center may cause more damage to the human resources in vicinity. According to a research given on Forbes portal [30], there were 10 billion devices interconnected till 2017 and the number could double in a single year. Improvement and extensions in scheduling tasks to the servers have to be added so that even a minimal reduction in execution time could lead to a much needed hoard of electrical power being saved.

### B. Measure of Energy Efficiency

Data center efficiency is measured in terms of Power Usage Effectiveness (PUE) which gives the ratio of total data center energy usage to the energy usage of IT equipment in the data center. This shows that a lower PUE value means reduced energy consumption in cooling and more in computing needs. It also means that the device consumes less energy with respect to the whole data center energy consumption with less greenhouse emissions. Environment Protection Agency's Energy Star program reported a PUE of 1.91 in 2009 after collecting information from 100 data centers [32]. Current PUE measurement for Google data centers is reported to be 1.12 making them among the most efficient in the world [27]. With lower PUE, cloud providers need to focus on QoS (Quality of Service) side by side in the rapidly increasing competitive market of cloud. According to the Green Grid [31], the most useful measurement point for measuring PUE is at the output of the Power Distribution Units (PDUs). This measurement should represent the total power delivered to the server racks in the data center.

The remainder of the paper is divided as follows. Section II gives a related study performed by known researchers in the stream

of energy consumption in cloud keeping different viewpoints for the reader. Section III introduces the implementation work on the mentioned topic. Section IV presents the simulation scenario for both, time-shared and space-shared task scheduling in the processing cores. Section V represents the conclusion and discussion.

## II. BACKGROUND STUDY

In [1], it was demonstrated that communication equipment form a larger than expected part of various located data centers and ignoring the power consumption of these equipment would not count as effective. Crucial factors which are associated to these components could be time delay in output and differences in bandwidth. Latency can be tolerated in video conferencing and bandwidth can be tolerated in scenario like cloud video streaming but only few scenarios exist where both can be tolerated. Network traffic is another rising situation where by traffic, it means the communication clash of data between resources with size expanding by 30 percent every year. The energy consumption in transport and switching can form a significant percentage of total energy consumption in cloud computing. Cloud computing can enable more energy-efficient use of computing power, especially when the user's predominant computing tasks are of low intensity or arise infrequently. However, that under some circumstances cloud computing can consume more energy than conventional computing on a local computer (PC). But multiple users are served as compared to computing done by a single PC. The broad point is that cloud computing offers significant energy savings through techniques like virtualization, consolidation of servers and advanced cooling systems.

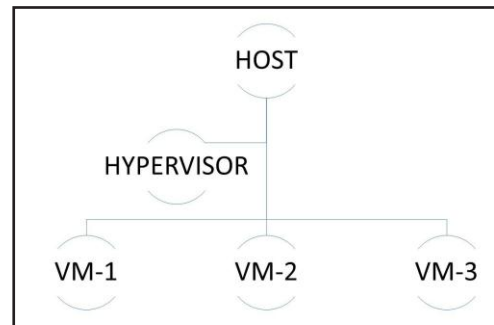


Fig. 2: Control of Hypervisor Over the Hosts to Manage the VMs

Going towards a lower level, data center consists of many server machines that are called hosts, these hosts further contain virtual machines. These machines are used in computing multiple tasks or cloudlets. A hypervisor (Fig. 2) is implemented on this level to control the VMs. A hypervisor is simply a virtual machine manager system that either works on the full host or comes in the form of a process, that depends whether it is a type-1 or type-2 hypervisor. For better reliability and performance, resources can be replicated [7] at redundant locations and using redundant infrastructures. To address exponential increase in

data traffic [28] and optimization of energy and bandwidth in data center systems, several data replication approaches have been proposed. New replicas need to be synchronized and any changes made at one of the sites need to be replicated at other locations. Energy is wasted in this process.

In the cloud computing approach multiple data center applications are hosted on a common set of servers. This allows for consolidation (Table I) of application workloads on a smaller number of servers that may be kept better utilized, as different workloads may have different resource utilization footprints and may further differ in their temporal variations. Consolidation thus allows amortizing the idle power costs more efficiently [18]. Henceforth proposed a novel techniques for power consolidation of VMs in data centers with the min-max and shares features which are inherent in virtualization technologies. These techniques provide a good power-performance trade-offs in modern data centers running multiple applications, wherein the number of resources allocated to a VM can be adjusted based on available resources, power costs, and application utilities. All major virtualization vendors provide parameters like min, max and shares. Setting a min for a VM confirms that it receives at least that amount of resources when

powered on and setting a max for a low priority application ensures that it does not use more resources, thus keeping them available for high-priority applications. Shares provide advice to the virtualization scheduler on the distribution of resources between contending VMs. Min, max and shares are useful only under high load conditions, but these conditions will occur very frequently in modern virtualized data centers. Some concerns and limitations which cannot be ignored in performing power consolidation while providing required performance are, Deciding which workloads should be combined on a common physical server. Resource usage, performance, and energy usage are not additive. Understanding the nature of their composition is thus critical to decide which workloads can be packed together. Secondly, there exists an optimal performance and energy point. Consolidation leads to performance degradation that causes the execution time to increase, eating into the energy savings from reduced idle energy. The main limitation is that performance degradation also occurs with consolidation because of internal conflicts among consolidated applications, for instance, cache contentions, conflicts at functional units of the CPU, disk scheduling conflicts, and disk write buffer conflicts.

TABLE I: COMPARISON TABLE OF TECHNIQUES DISCUSSED

| Techniques                                                                        | Description                                                                                                                               | Findings                                                                                                                                          |
|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------|
| Decentralized Content Aware Load Balancing                                        | Uses a unique and special property of requests and computing nodes to help scheduler to decide the best node for processing the requests. | Improves the searching performance and hence reducing idle time of nodes to save power.                                                           |
| Central load balancing policies for VMs                                           | Uses global state information to make load balancing decisions.                                                                           | Upto 20 percent improvement in performance. Also, does not consider fault tolerance and power management. Does not work according to the changes. |
| Static Data Replication (like RFH)                                                | Host node and replicas are predefined                                                                                                     | In recent popularity of data but the cost is low. Depends on usage pattern of data along.                                                         |
| Dynamic Data Replication (like RTRM and D2RS)                                     | Intensifies number of replicas according to average response time to user requests.                                                       | With storage pattern for storing that data. Knowledge of these two is important.                                                                  |
| Server consolidation using Multi-dimensional Bin-packing (or min, max and shares) | Server like resources act as bins and their capacity gives a decision to whether turn them on or off.                                     | Experiments were conducted beforehand with given capacities of resources.                                                                         |
| Two Threshold policy                                                              | Keeps record of the two threshold values for CPU utilization for reallocation decisions.                                                  | Only CPU is treated as a resource leaving other essential resources.                                                                              |

Another research showed that not only server based usage of resources indicate energy dissipation, network components like switches and bandwidth utilized in the channels associated with those switches are sufficient to give a sight of power consumption. The objective of [7] paper was to minimize utilization of network bandwidth along with communication delays encountered in a data center network. Bandwidth available per server plays an important role to reduce delays in transfer, even the widely used three-tier architecture mentioned earlier accounts for an effective number of access, core and aggregation switches along with number of servers kept per rack. According to the study conducted, two communication links are mainly associated with a data center, namely uplink flow, that direct from servers to core switches and downlink flow, directed from core switches to servers. These two types of flows act as a channel between the database objects and servers requesting those objects at the time of execution. The effect of applying DVFS is considered as limited as it is applied only to the CPU, leaving behind system peripherals like buses, memory and disks.

Virtual Machine placement is a complex task to handle when enormous amount of users are considered at the same time, discussing this [17], proposed an energy-efficient resource management strategy that considers three ultimate goals. The first goal was to reallocate virtual machines according to current CPU, RAM and network bandwidth at each time frame. Second goal was to consider optimization of network communication and protocols for intercommunication of those virtual machines. This also reduces data transfer overhead from different types of network devices. The final goal was to provide efficient cooling and thermal optimization of nodes in order of their current states. The research provided a tiered architecture for resource management that included the dispatchers for distributing the requests, the local managers that formed a part of the Virtual Machine Monitor (VMM) taking account of the current utilization of the node's resources and their thermal state, and other tier was of the global managers that finally take the reallocation decision after applying a heuristic approach for bin-packing where bins are the nodes and items are the VMs. For removing a Single Point of Failure (SPF), global managers are not included inside a physical node but are attached to a set of nodes for processing the request of local managers.

### III. IMPLEMENTATION DETAILS

The cloudsims toolkit [25] provides various program classes for handling different assets of a data center and to simulate the communication between the user and the data center. The user is also referred to as broker in the toolkit. There is a separate class for data center, data center broker and cloud information service. The simulation presented in this paper gives a definite comparison between the use of time-shared and space-shared policies applied inside a virtual machine to arrange tasks given by the user to complete. The task may include entering data, processing data, and access or storage task. These tasks are referred to as cloudlets in the toolkit. Each task is of random

length known to the server earlier but the algorithm applied here knows the length of each task that arrives at cloud information service which acts as a gateway to the data center.

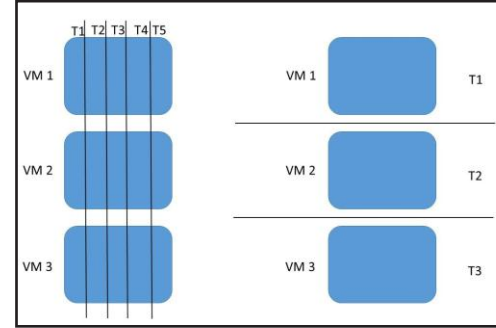


Fig. 3: Division of Cores According to Time-Shared and Space-Shared Policies in CloudSim. T1, T2, T3 are Tasks Assigned to the VMs

Assuming the basic routine of energy that is used in processing each cloudlet, it depends on time of execution of the cloudlets if each cloudlet consumes the same amount of power. We assume the network factors like bandwidth and RAM to be constant in this case to give an impact of time on the total energy. The energy consumed by the processors depends on the number of cloudlets given for processing.

Fig. 4 illustrates the working of time-shared and space-shared policies that help in assigning the cloudlets to single or multiple VMs. In time-shared policy, multiple virtual machines complete a single task by sharing the execution time of the task. Vigorous context switching is needed to accomplish this. Time-sharing can also be compared to pre-emptive scheduling done in operating systems where a specific quantum of time is provided to a process so that each process could run. Space-shared policy ensures the queuing of tasks on a particular virtual machine so that no task shares the machine with another task at the same time. Each CPU core is given a single VM and each VM is given a single task.

TABLE II: CONSTANT PARAMETERS FOR THE SIMULATION

| Parameter       | Size (MB) | RAM (MB) | MIPS | Bandwidth |
|-----------------|-----------|----------|------|-----------|
| Virtual Machine | 512       | 1000     | 1000 | 1000      |
| Cloudlet        | 512       | 1000     | 1000 | 1000      |

In Table II, the configuration of a single instance of a VM and a cloudlet in CloudSim is represented which remained constant throughout. The size and RAM allocated to the tasks in total are measured in Mega Bytes (MB). The equation used to observe the result is

$$E = \oint_n P \cdot dt$$

where,  $E$  is energy,  $n$  is limit of integral,  $P$  is power,  $dt$  is change in time.

Where,  $n$  varies with time. Here,  $n$  is the number of cloudlets,  $P$  is the constant power assumed, and  $t$  is the processing time taken by the cloudlet. The total execution time for all cloudlets is calculated by the equation where, the number of tasks vary from  $i=0$  to  $i=n$ , and each  $C_i(t)$  represents the time of execution of each cloudlet that is also a function of resistance created in the communication network which is out of scope of this paper. For each simulation the variable value for the number of virtual machines could be varied. Main results were accounted for 5 virtual machines with 40 cloudlets for the first time and 60 cloudlets for the second time. These numbers gave the result fast enough without any lag and delay.

$$T = \sum_{i=0}^n C_i(t)$$

TABLE III: RESULT AFTER APPLYING SPACE-SHARED CLOUDLET SCHEDULING

| Cloudlets | VMs | Exec. Time (ms) | Process. Time (ms) |
|-----------|-----|-----------------|--------------------|
| 40        | 5   | 16.4            | 79.18              |
| 60        | 5   | 23.96           | 117.57             |

TABLE IV: RESULT AFTER APPLYING SPACE-SHARED CLOUDLET SCHEDULING

| Cloudlets | VMs | Execution Time (ms) |
|-----------|-----|---------------------|
| 40        | 5   | 16.72               |
| 60        | 5   | 25.94               |

In the simulation, each cloudlet is assumed to be of different length and a normal ascending ordering of the cloudlets is done lengthwise. This is done for improving the overall effective ratio of number of tasks completed per unit time because this ratio will be greater at the initial phase of the process.

Each VM gets to serve the cloudlets. In time-shared policy, the VMs work together to process a request and in space-shared policy, a queue is maintained to keep the tasks waiting till the earlier tasks are processed. Simulation for no-policy applied execution is also observed.

#### IV. SIMULATION RESULTS

Table III represents the application of space-shared policy after sorting the cloudlets according to length, the result showed a more drastic change when number of cloudlets were increased from 40 to 60, and execution time was reduced at a measurable scale after applying the time-shared policy and keeping all other parameters same.

We know that the execution time for time-shared policy should be less than that of space-shared policy. It was noted that the change in execution time of each cloudlet was different because

of its length but this change could not be noted in time-shared policy. The execution time for each cloudlet is shown in terms of total execution time given in Table III and Table IV. Time-shared policy needs all processing elements or cores to work and be used. But using all processing elements with normal utilization could only be done in space-shared policy because there would be more load on the cores when time-shared policy is used. This load is caused because of the vigorous context switching mentioned earlier. Time-shared policy keeps the cores busy even when a single task is given for the machine to compute. Queuing of these tasks in space-shared policy help in taking control of their execution.

Also, as per the cloudbus portal, the time-shared policy suffers from 10 percent performance degradation due to VM migration. Time-shared policy has much more better execution time than space-shared but it is vulnerable to higher risk for data loss due to server crashes. Space-shared policy is better when applied to larger number of cloudlets so that each processing element is given a chance to execute and become free for execution of the next task. Energy of the system is observed in terms of time, considering space-shared policy for greater amount of tasks. Different sorting algorithms can be applied to improve the performance of the system in Data Center Broker (data center broker) class in the toolkit. But in this paper simple sorting is applied.

#### V. CONCLUSION

During the simulation, a keen observation was made about the time of execution of each cloudlet and the change in this time was not constant between any of two successive cloudlets except in time-shared policy of scheduling. The choice between the two techniques, time-shared and space-shared depends on the user according to the scenario presented in the simulation results. Main concern that affects the execution is the cloudlet scheduling algorithm applied that performs the task of sending appropriate size cloudlet to an appropriate core for processing. Here, a simple algorithm of sorting is used to assign the cores to the least sized cloudlet first. These algorithms will be further improved in future research work, thus decreasing the execution time and hence, the reduction in energy usage.

#### REFERENCES

- [1] J. Baliga, R. W. A. Ayre, K. Hinton, and R. S. Tucker, "Green cloud computing: Balancing energy in processing, storage, and transport," *Proceedings of the IEEE*, vol. 99, no. 1, pp. 149-167, 2011.
- [2] D. S. Markovic, D. Zivkovic, I. Branovic, R. Popovic, and D. Cvetkovic, "Smart power grid and cloud computing," *Renewable and Sustainable Energy Reviews*, vol. 24, pp. 566-577, 2013.
- [3] N. J. Kansal, and I. Chana, "Cloud load balancing techniques: A step towards green computing," *IJCSI*

- International Journal of Computer Science*, vol. 9, no. 1, pp. 238-246, 2012.
- [4] R. Mata-Toledo, and P. Gupta, "Green data center: How green can we perform?," *Journal of Technology Research*, 2011.
  - [5] A. Beloglazov, R. Buyya, Y. C. Lee, and A. Zomaya, "A taxonomy and survey of energy-efficient data centers and cloud computing systems," *Advances in Computers*, vol. 82, pp. 47-111, 2011.
  - [6] A. Banerjee, P. Agrawal, and N. Ch. S. N. Iyengar, "Energy efficiency model for cloud computing," *International Journal of Energy, Information and Communications*, vol. 4, no. 6, pp. 29-42, 2013.
  - [7] D. Boru, D. Kliazovich, F. Granelli, P. Bouvry, and A. Y. Zomaya, "Energy-efficient data replication in cloud computing datacenters," *Cluster Computing*, vol. 18, no. 1, pp. 385-402, 2015.
  - [8] X. Dong, T. El-Gorashi, and J. M. H. Elmirghani, "Green IP over WDM networks with data centers," *Journal of Lightwave Technology*, vol. 29, no. 12, pp. 1861-1880, 2011.
  - [9] M. Lin, A. Wierman, L. L. H. Andrew, and E. Thereska, "Dynamic right-sizing for power-proportional data centers," *IEEE/ACM Transactions on Networking*, vol. 21, no. 5, pp. 1378-1391, 2013.
  - [10] S. Ghemawat, H. Gobioff, and S. T. Leung, "The Google file system," *ACM SIGOPS Operating Systems Review*, vol. 37, no. 5, p. 29, 2003.
  - [11] B. A. Milani, and N. J. Navimipour, "A comprehensive review of the data replication techniques in the cloud environments: Major trends and future directions," *Journal of Network and Computer Applications*, vol. 64, pp. 229-238, 2016.
  - [12] Y. Qu, and N. Xiong, "RFH: A resilient, fault-tolerant and high-efficient replication algorithm for distributed cloud storage," In *2012 41<sup>st</sup> International Conference on Parallel Processing (ICPP)*, IEEE, 2012.
  - [13] X. Bai, H. Jin, X. Liao, and Z. Shao, "RTRM: A response time-based replica management strategy for cloud storage system," *International Conference on Grid and Pervasive Computing*, Springer, Berlin, Heidelberg, 2013.
  - [14] D. W. Sun, G. R. Chang, S. Gao, L. Z. Jin, and X. W. Wang, "Modeling a dynamic data replication strategy to increase system availability in cloud computing environments," *Journal of Computer Science and Technology*, vol. 27, no. 2, pp. 256-272, 2012.
  - [15] T. Hamrouni, S. Slimani, and F. B. Charrada, "A survey of dynamic replication and replica selection strategies based on data mining techniques in data grids," *Engineering Applications of Artificial Intelligence*, vol. 48, pp. 140-158, 2016.
  - [16] S. Srikantaiah, A. Kansal, and F. Zhao, "Energy-aware consolidation for cloud computing," *Proceedings of the 2008 Conference on Power Aware Computing and Systems*, vol. 12, no. 1, p. 10, 2008.
  - [17] A. Beloglazov, and R. Buyya, "Energy efficient resource management in virtualized cloud data centers," *Proceedings of the 2010 10<sup>th</sup> IEEE/ACM International Conference on Cluster, Cloud and Grid Computing*, IEEE Computer Society, 2010.
  - [18] M. Cardosa, M. R. Korupolu, and A. Singh, "Shares and utilities based power consolidation in virtualized server environments," *2009 IFIP/IEEE International Symposium on Integrated Network Management*, IEEE, 2009.
  - [19] A. Beloglazov, J. Abawajy, and R. Buyya, "Energy-aware resource allocation heuristics for efficient management of data centers for cloud computing," *Future Generation Computer Systems*, vol. 28, no. 5, pp. 755-768, 2012.
  - [20] R. Raghavendra, P. Ranganathan, V. Talwar, Z. Wang, and X. Zhu, "No power struggles: Coordinated multi-level power management for the data center," *ACM SIGARCH Computer Architecture News*, ACM, vol. 36, no. 1, 2008.
  - [21] J. Shuja, K. Bilal, S. A. Madani, M. Othman, R. Ranjan, P. Balaji, and S. U. Khan, "Survey of techniques and architectures for designing energy-efficient data centers," *IEEE Systems Journal*, vol. 10, no. 2, pp. 507-519, 2016.
  - [22] N. Liu, Z. Dong, and R. Rojas-Cessa, "Task scheduling and server provisioning for energy-efficient cloud-computing data centers," *2013 IEEE 33<sup>rd</sup> International Conference on Distributed Computing Systems Workshops (ICDCSW)*, IEEE, 2013.
  - [23] N. Liu, Z. Dong, and R. Rojas-Cessa, "Task and server assignment for reduction of energy consumption in data centers," *2012 11<sup>th</sup> IEEE International Symposium on Network Computing and Applications (NCA)*, IEEE, 2012.
  - [24] D. Bouley, "Estimating a data centers electrical carbon footprint," *Schneider Electric White Paper Library*, 2011.
  - [25] R. N. Calheiros, R. Ranjan, A. Beloglazov, C. A. F. D. Rose, and R. Buyya, "CloudSim: A toolkit for modeling and simulation of cloud computing environments and evaluation of resource provisioning algorithms," *Software: Practice and Experience*, vol. 41, no. 1, pp. 23-50, 2011.
  - [26] M. Erol-Kantarci, and H. T. Mouftah, "Energy-efficient information and communication infrastructures in the

- smart grid: A survey on interactions and open issues,” *IEEE Communications Surveys Tutorials*, vol. 17, no. 1, pp. 179-197, 2015.
- [27] <https://www.google.com/about/datacenters/efficiency/internal/>
- [28] Cisco, Cisco Visual Networking Index: Forecast and Methodology, 2011-2016, White Paper, May 2012. Available: [http://www.cisco.com/en/US/solutions/collateral/ns341/ns525/ns537/ns705/ns827/white\\_paper\\_c11-481360.pdf](http://www.cisco.com/en/US/solutions/collateral/ns341/ns525/ns537/ns705/ns827/white_paper_c11-481360.pdf)
- [29] P. Mell, and T. Grance, “The NIST definition of cloud computing,” 2011.
- [30] <https://www.forbes.com/sites/forbestechcouncil/2017/12/15/why-energy-is-a-big-and-rapidly-growing-problem-for-data-centers/7fa591b05a30>
- [31] Green Grid, “The Green Grid data center power efficiency metrics: PUE and DCiE,” Green Grid Report, 2007.
- [32] <http://www.datacenterknowledge.com/archives/2011/05/10/uptime-institute-the-average-pue-is-1-8>