

 Open access • Journal Article • DOI:10.1007/S10586-014-0404-X

Energy-efficient data replication in cloud computing datacenters — [Source link](#)

Dejene Boru, Dzmitry Kliazovich, Fabrizio Granelli, Pascal Bouvry ...+1 more authors

Institutions: University of Luxembourg, University of Trento, University of Sydney

Published on: 01 Mar 2015 - IEEE International Conference on Cloud Computing Technology and Science

Topics: Cloud computing, Replication (computing), Quality of service, Bandwidth (computing) and Bottleneck

Related papers:

- [CloudSim: a toolkit for modeling and simulation of cloud computing environments and evaluation of resource provisioning algorithms](#)
- [Energy-aware resource allocation heuristics for efficient management of data centers for Cloud computing](#)
- [MORM: A Multi-objective Optimized Replication Management strategy for cloud storage cluster](#)
- [A comprehensive review of the data replication techniques in the cloud environments](#)
- [Modeling a Dynamic Data Replication Strategy to Increase System Availability in Cloud Computing Environments](#)

Share this paper:    

View more about this paper here: <https://typeset.io/papers/energy-efficient-data-replication-in-cloud-computing-1ijrbbigx>

Energy-Efficient Data Replication in Cloud Computing Datacenters

Dejene Boru¹, Dzmitry Kliazovich², Fabrizio Granelli³, Pascal Bouvry², Albert Y. Zomaya⁴

¹ CREATE-NET

Via alla Cascata 56/D, 38123, Trento, Italy
dejene.oljira@create-net.org

² University of Luxembourg

6 rue Coudenhove Kalergi, Luxembourg
{dzmitry.kliazovich, pascal.bouvry}@uni.lu

³ DISI - University of Trento

Via Sommarive 14, Trento, Italy
granelli@disi.unitn.it

⁴ School of Information Technologies

University of Sydney, Australia
albert.zomaya@sydney.edu.au

Abstract- Cloud computing is an emerging paradigm that provides computing resources as a service over a network. Communication resources often become a bottleneck in service provisioning for many cloud applications. Therefore, data replication, which brings data (e.g., databases) closer to data consumers (e.g., cloud applications), is seen as a promising solution. It allows minimizing network delays and bandwidth usage. In this paper we study data replication in cloud computing data centers. Unlike other approaches available in the literature, we consider both energy efficiency and bandwidth consumption of the system, in addition to the improved Quality of Service (QoS) as a result of the reduced communication delays. The evaluation results obtained during extensive simulations help to unveil performance and energy efficiency tradeoffs and guide the design of future data replication solutions.

Keywords: Cloud computing, data replication, energy efficiency

I. INTRODUCTION

Cloud computing is an emerging technology that attracts ICT service providers offering tremendous opportunities for online distribution of services. It offers computing as a utility, sharing resources of scalable data centers. End users can benefit from the convenience of accessing data and services globally, centrally managed backups, high computational capacity, and flexible billing strategies. Cloud computing is also ecologically friendly. It benefits from the efficient utilization of servers, data center power planning, large scale virtualization, and optimized software stacks. However, electricity consumed by the data centers is still in the order of thousands of megawatts. Global power demand of data centers in 2005 was equivalent to seventeen 1,000 MW power plants, which corresponds to 1% of the worldwide electricity use [1]. In 2010, this number was already topping 1.5% of global electricity consumption [2]. The growth of Internet services at an unprecedented rate requires the development of novel optimization techniques at all levels to cope with escalation in energy consumption, which in place would reduce operational costs and carbon emissions.

Data centers typically overprovision computing, storage, power distribution, and cooling infrastructures to ensure high levels of reliability [3]. The cooling and power distribution systems consume around 45% and 15% of the total energy

respectively, while leaving roughly 40% to the IT equipment [4]. These 40% are shared between computing servers and networking equipment. Depending on the data center load level, the communication network can consume between 30 and 50% of the total power of the IT equipment [5], [6].

There are two main approaches for making data center consume less energy: shutting the components down or scaling down their performance. Both approaches are applicable to computing servers [7], [8] and network switches [5].

The performance of cloud computing applications such as gaming, voice and video conferencing, online office, storage, backup, social networking, depends largely on the availability of high-performance communication resources and network efficiency [10]. For better reliability and high performance low latency service provisioning, the data resources can be brought closer (replicated) to the physical infrastructure, where the cloud applications are executed. Therefore, a large number of replication strategies for data centers have been proposed [3], [11], [12], [13], [14]. These strategies optimize system bandwidth and data availability providing replication strategies between geographically distributed data centers. However, none of them focuses on energy efficiency and replication techniques inside data centers.

To address these gaps, we propose a data replication technique for cloud computing data centers which optimizes energy consumption, network bandwidth, and communication delays which can be applied in both the geographically distributed data centers as well as inside each individual data center. Specifically, our contributions can be summarized as follows.

- Development of a data replication approach for joint optimization of energy consumption and bandwidth capacity of data centers.
- Optimization of communication delays to ensure the quality of user experience with cloud applications.
- Performance evaluation of the developed replication strategy using packet-level cloud computing simulator GreenCloud [15].
- Analysis of the tradeoff between performance, serviceability, reliability and energy consumption.

The rest of the paper is organized as follows: Section II highlights relevant works related to our contribution. In Section III we introduce the energy efficient replication solution. Section IV presents evaluation details of our proposed solution. Finally, conclusions and future work are outlined in Section V.

II. RELATED WORKS

A. Energy efficiency

At the component level, there are two main alternatives for making data center consume less energy: (a) shutting the hardware components down or (b) scaling down their performance. These methods are applicable to both: computing servers and network switches.

When applied to the servers, the former method is commonly referred to as Dynamic Power Management (DPM) [7]. DPM results in most of the savings. It is most efficient when combined with the workload consolidation scheduler – the policy which allows maximizing the number of idle servers that can be put into a sleep mode, as the average load of cloud computing systems often stays below 30% [7]. The second method corresponds to the Dynamic Voltage and Frequency Scaling (DVFS) technology [8]. DVFS exploits the relation between power consumption P , supplied voltage V , and operating frequency f and is given with the expression:

$$P = V^2 * f. \quad (1)$$

Reducing voltage or frequency reduces the power consumption. The effect of DVFS is limited as power reduction applies only to the CPU, while system bus, memory, disks as well as peripheral devices continue to consume at their peak rates.

Similar to the computing servers, most of the energy-efficient solutions for communication equipment depend on (a) downgrading the operating frequency (or transmission rate) or (b) powering down the entire device or its components in order to conserve energy. Power-aware networks were first studied by Shang et al. [5]. In 2003, the first work that proposed a power-aware interconnection network utilized Dynamic Voltage Scaling (DVS) links [5]. After that, DVS technology was combined with Dynamic Network Shutdown (DNS) to further optimize energy consumption.

Another technology, which affects energy consumption indirectly and is currently widely adapted, is virtualization [16]. Virtualization allows sharing of a single physical server among multiple virtual machines (VM), where each VM can serve different applications. The resources of servers can be dynamically provisioned to a VM based on the application requirements. When applied to computing servers virtualization allows sharing of computing resources, such as processors or memory, among different applications resulting in reduced power consumption of physical servers [17].

In networking domain, the virtualization enables implementation of logically different addressing and forwarding mechanisms, which may not necessarily have the goal of energy efficiency.

B. Data Replication

The emergence of cloud computing enabled the deployment of immense IT services that are built on top of geographically

distributed platforms and offered globally. For better reliability and performance, resources are replicated at the redundant locations and using redundant infrastructures. To address an exponential increase in Internet data traffic and optimize energy and bandwidth in datacenter systems, several data replication approaches have been proposed.

The approach presented in [11] involves turning off underutilized storage servers to minimize energy consumption. To guarantee availability, one replica is kept for every data object. A mixture of two data layout policies is proposed: the sequential data layout policy which allows shutting down all replica servers but one and the random data layout policy which allows no more than $k - 1$ servers to be turned off. Here k is the number of replica servers that hold one particular data object. The latter policy has an advantage of having multiple nodes participating in the recovery process, while the former policy is energy efficient.

A dynamic data replication approach in cluster of data grids is proposed in [3]. A policy maker manages replicas. It collects information from cluster heads and determines file popularity based on the access frequency. The optimal number of replicas is computed based on the relative access rate of all the files in the system. This approach is centralized and therefore exposed to a single point of failure. Moreover, energy efficiency is not taken into account. The replication strategy across datacenters to reduce energy consumption of backbone network is proposed in [12]. Optimal replication site is selected based on the data center traffic demands and popularity of a data object using linear programming. This work focuses on replication strategies between data centers, but replication inside data centers is not considered.

In [13], data replication across datacenters with the objective of reducing access delay is proposed. The Optimal replication site is selected based on the access history of the data. A weighted k-means clustering of user locations is used to determine replica site location. The replica is deployed closer to the central part of each cluster. A cost-based data replication in cloud datacenter is proposed in [14]. This approach analyzes data storage failures and data loss probability that are in the direct relationship and builds a reliability model. Then, replica creation time is determined by solving reliability function.

The approach presented in this paper is different from all replication approaches discussed above (a) by the scope, which implements data replication both within a data center as well as between geographically distributed data centers, (b) by the optimization target, which takes into account system energy consumption, network bandwidth, and communication delays.

III. ENERGY-EFFICIENT REPLICATION

In this paper we assume multiple cloud computing datacenters geographically distributed across the globe (see Fig. 1). Each datacenter has a three tier topology. Its interconnection network comprises of the core, aggregation, and access layers. The core layer provides packet switching backplane for all the flows going in and out of the datacenter. The aggregation layer integrates connections and traffic flows from multiple racks. The access layer is where computing servers are arranged into racks.

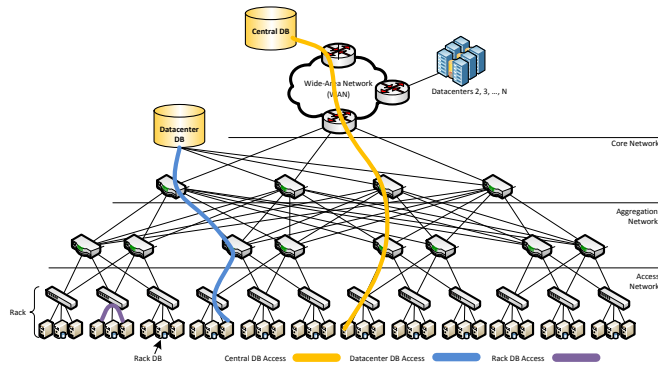


Fig. 1. Three-tier cloud computing data center architecture.

A central database (Central DB), located in the wide-area network, hosts all the data required by the cloud applications. To speed up the access and reduce latency, each data center hosts a local database, called datacenter database (Datacenter DB). It is used to replicate the most frequently used data items from the central database. Each rack hosts at least one server capable of running local rack-level database (Rack DB), which is used for replication of data from the datacenter database.

Any database request generated by the cloud applications running at the computing servers is initially directed to a rack-level database server. This server either replies with the requested data or forwards the request to the datacenter database. In a similar fashion, the datacenter database either satisfies the request or forwards it up to the central database.

When data is accessed, the information about requesting server, the rack, and the datacenter is stored. In addition, the statistics showing the number of accesses and updates are obtained for each data item. The access rate (or popularity) is measured as the number of access events in a given period of time. The popularity is not constant. It varies in relation to the types of the stored data. Typically, a newly created data have the highest demand. Then, the access rate decays over time. For example, a newly posted YouTube video attracts most of the visitors. However, as the time passes its popularity and audience start to decay [18].

A module called replica manager is located at the central database. It periodically analyzes data access statistics to identify which data items are the most suitable for replication and at which replication sites. The availability of the access and update statistics makes it possible to project data center bandwidth usage and energy consumption.

From the network bandwidth perspective, the availability of per-server bandwidth capacity is one of the core requirements affecting the design of modern data centers. The most widely used three-tier fat tree topology (see Fig. 1) has strict limits in the number of hosted core, aggregation, and access switches as well as the number of servers per rack. For example, a rack switch serving 48 servers connected with 1 Gb/s link has only two 10 Gb/s links in the uplink to aggregation switches. This corresponds to oversubscription ratio of 2.4 which limits per-server available bandwidth to 416 Mb/s. Further bandwidth multiplexing occurs at the aggregation layer with access and higher layer core switches. An aggregation switch offers 12 ports for access layer connections and is connected to all the core layer switches. For three tier architecture with 8-way

Equal Cost Multipath Routing (ECMP) [19], the aggregation layer oversubscription ratio is 1.5. This further reduces the actual per server bandwidth down to 277 Mbps.

Most of the cloud applications, such as online office and social networking, rely on tight interaction with databases. Data queries can be fulfilled either locally or from a remote site. To ensure data availability and reduce access delays data replication can be used.

Fig. 2 presents a timeline related to a workload execution in data center. The starting point is when the user request arrives to the datacenter gateway. After being scheduled it is forwarded through data center network to the selected computing server for execution. At the server, if needed, the workload requests data required for its execution by sending a database request and waiting for a reply. The database querying delay corresponds to the round-trip time which varies depending on the database location. As soon as the reply is received, the workload execution is started. Some of the executed workloads modify the received data item and send the update back to the database at the end of their execution.

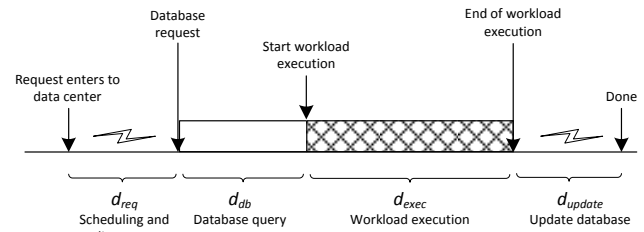


Fig. 2. Workload execution timeline.

According to the model presented in Fig. 2, all data transmissions in data center can be broadly divided in to the uplink and downlink types. The uplink flows are those directed from the computing servers towards the core switches, while the downlink flows are those from the core switches to the computing servers.

In the uplink, the available network bandwidth is used for propagating database requests and when applications need to update the modified data item. In the downlink, the bandwidth is used for delivering the workload descriptions to the servers for execution, receiving database objects, and propagating updates between database replicas.

The main goal of the proposed replication strategy is to improve system performance while minimizing the energy consumption and bandwidth usage. In particular, we target minimization of the datacenter energy consumption and the maximization of the residual bandwidth left unused in the downlink as well as in the uplink.

All databases maintain and exchange access statistics recording data access and update rates for each data item. In addition, congestion levels in different parts of the datacenter network are monitored. In a fat tree topology the bandwidth is concentrated towards the root of the tree. Therefore, data replication at lower levels may significantly increase system performance.

IV. PERFORMANCE EVALUATION

This section presents performance evaluation of the proposed replication strategy. The main performance indicators are: data center energy consumption, available network bandwidth, and communication delay. The following subsections present details about power models, describe simulation scenario, and analyze the obtained results.

A. Power Consumption Models for Servers and Network Switches

The power model followed by server components is dependent on the server state and its CPU utilization. As reported in [20], an idle server consumes about two-thirds of its peak power consumption. This is due to the fact that servers must manage memory modules, disks, I/O resources, and other peripherals in an acceptable state. Then, power consumption linearly increases with the level of CPU load. In this work we consider a more detailed model [20], which is largely based on experimental results and suggests non-linear distribution of the workload-dependent power consumption. This non-linear power consumption of the server can be computed as follows:

$$P_s(l) = P_{fixed} + \frac{(P_{peak} - P_{fixed})}{2} (1 + l - e^{-\frac{l}{a}}), \quad (2)$$

where P_{fixed} is an idle power consumption, P_{peak} power consumed at the peak load, l is a server load, and a is a scaling coefficient which is typically in the range [0.2, 0.5] and corresponds to the utilization level at which the server attains asymptotic power consumption.

Network switches are physical layer devices that include such hardware components as port transceivers, line cards, and chassis. Several studies have characterized energy consumption of network switches [22]. The suggested energy model [23] combines static and dynamic parts. The static part, related to the power consumed by chassis and line cards, is independent of the traffic load. On the contrary, the dynamic part, which characterizes the consumption of network ports, depends on the traffic activity. Hence, energy consumption of a network switch can be given as the following:

$$P_{switch} = P_{chassis} + n_c * P_{linecard} + \sum_{r=1}^R n_p^r * P_p^r * u_p^r, \quad (3)$$

where $P_{chassis}$ is a power related to switch chassis, $P_{linecard}$ is the power consumed by a single line card, n_c is number of line cards plugged into switch, P_p^r is power drawn by a port running at rate r , n_p^r is number of ports operating at rate r and $u_p^r \in [0,1]$ is a port utilization which can be defined as follows:

$$u_p = \frac{1}{T} \int_t^{t+T} \frac{B_p(t)}{C_p} dt = \frac{1}{T * C_p} \int_t^{t+T} B_p(t) dt, \quad (4)$$

where $B_p(t)$ is an instantaneous throughput at the time t , C_p is the link capacity, and T is a measurement interval.

Table I presents the values used to feed power presented above. The server peak energy consumption of 301 W includes 130 W allocated for a peak CPU consumption [24] and 171 W consumed by other devices like memory, disks, peripheral slots, mother board, fan, and power supply unit [20]. As the

only component which scales with the load is the CPU power, the consumption of an idle server is bounded by 198 W.

TABLE I. POWER CONSUMPTION OF DATACENTER HARDWARE

Parameter	Power Consumption [W]		
	Chassis	Line cards	Port
Gateway, core, aggregation switches	1558	1212	27
Access switches	146	-	0.42
Computing server	301		

Energy consumption of network switches is almost constant for different transmission rates as 85-97% of the power is consumed by switches' chassis and line cards and only a small portion of 3-15% is consumed by their port transceivers. The values for power consumption are derived from [25].

B. Simulation Scenario

For performance evaluation purposes we developed the GreenCloud simulator [15] that was extended with a set of classes providing data replication functionalities. GreenCloud is a cloud computing simulator which captures data center communication processes at the packet level. It is based on Ns2 simulation platform [17] and allows capturing realistic TCP/IP behavior in a large variety of network scenarios. GreenCloud offers a detailed fine-grained modeling of the energy consumed by the data center hardware, such as servers, switches, and communication links, and implements a number of network-aware resource allocation and scheduling solutions [26], [21]. The performance evaluation considers only single datacenter with a holistic model for resource representation [9]. However, the obtained results can be easily extrapolated to the scenario with multiple datacenters.

Table II summarizes simulation setup parameters. The simulated data center is comprised of 1024 servers arranged into 32 racks which are interconnected by 4 core and 8 aggregation switches. The network links connecting the core and aggregation switches as well as the aggregation and access switches are 10 Gb/s. The bandwidth of the access links connecting computing servers to the top-of-rack switches is 1 Gb/s. The propagation delay of these links is set to 3.3 μ s. There is only one entry point to the datacenter through a gateway switch which is connected to all the core layer switches with 100 Gb/s, 50 ms links.

TABLE II. DATACENTER TOPOLOGY

Parameter	Value
Gateway nodes	1
Core switches	4
Aggregation switches	8
Access (rack) switches	32
Computing servers	1024
Gateway link	100 Gb/s, 50 ms
Core network link	10 Gb/s, 3.3 μ s
Aggregation network link	10 Gb/s, 3.3 μ s
Access network link	1 Gb/s, 3.3 μ s

The workload generation events are exponentially distributed in time to mimic typical process of user arrival. As soon as a scheduling decision is taken for a newly arrived workload it is sent over the data center network to the selected server for execution. The workload execution and data querying follow the timeline diagram presented in Fig. 2. The

size of a workload description and database queries is limited to 1500 bytes and fits into a single Ethernet packet. The size of data items, data access and update rates, as well as the replication threshold varies in different simulation runs. The duration of each simulation run is 60 minutes. DNS power saving scheme is enabled for both servers and switches in all simulation scenarios.

The following subsections report the effect of data size and the update rate on energy consumption, network bandwidth and access delay characteristics of the system.

C. Simulation Results

Fig. 3 presents system bandwidth requirements in the downlink, when no database updates are performed. Being proportional to both the size of a data item and the update rate, the bandwidth consumption grows fast and easily overcomes the capacities of different segments of the datacenter network requiring replication. The availability of only 100 Gb/s at the gateway link would trigger replication even for the small data items of less than 12 MB (or 8 Ethernet packets) for an access rate of 1 Hz. Thus, replication from Central DB to the Datacenter DB would be required to avoid bottleneck. The bandwidth provided by the core network of 320 Gb/s will be exceeded with data items larger than 40 MB for the access rate of 1 Hz. Similarly, the bandwidth of the aggregation network of 640 Gb/s will be exceeded after 78 MB is reached which will demand an additional data replication from Datacenter DB to Rack DBs. Finally, data sizes greater than 125 MB will cause traffic congestion at the largest in the system access segment of the network clearly identifying the limit.

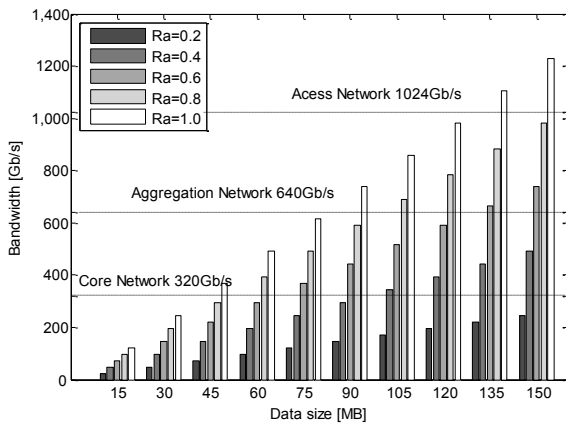


Fig. 3. Downlink bandwidth demand.

Fig. 4 presents the measurements of energy consumption of computing servers for data item sizes varied from 10 MB to 40 MB. Each server accesses one data item every 0.3 seconds and sends no updates back to the database. There are two trends that can be observed from the obtained results. The first trend is that energy consumption increases with the increase in data size. The second is that energy consumption decreases as data becomes available at closer to the computing servers locations. The reason is that the communication delay is included into the execution time of the cloud application (see Fig. 2), which prevents servers to enter into the sleep mode. These delays

become large with the increase of the data item size, but can be reduced by shortening the round-trip times to the database.

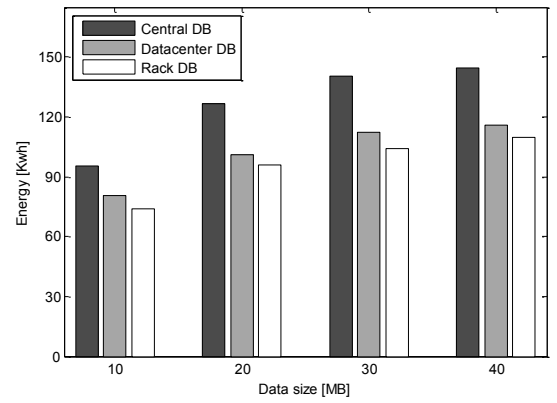


Fig. 4. Energy consumption of servers.

Fig. 5 presents energy consumption of network switches for the scenario with fixed data size of 6 MB and access rate of 0.3 Hz, but a variable update rate. As expected, it increases with the increase of the update rate due to longer awake periods. Switches at all layers are involved into forwarding database update traffic. In the uplink, they forward replica updates sent from the servers to the Central DB. In the downlink, database updates from Central DB to Datacenter DB and from Datacenter DB to Rack DBs are propagated. In the case of Datacenter DB replication (see Fig. 5 (b)), only the gateway and core switches are involved into update flow forwarding. In the case of Rack DB replication (see Fig. 5 (c)), both core and aggregation networks carry database updates for 32 Rack DBs. The access switches serve both data traffic and the database updates, which justify their higher energy consumption values.

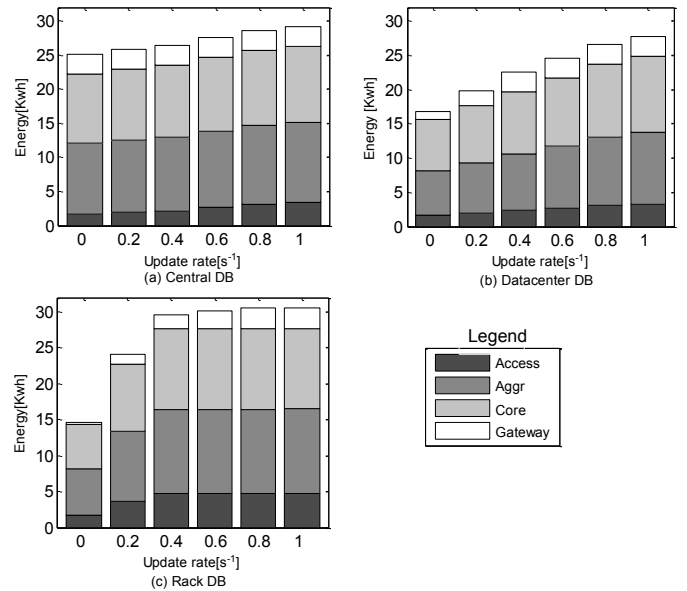


Fig. 5. Energy consumption of network switches.

Fig. 6 reports data access delays measured as an average time elapsed since sending data request and the time the requested data are received. As expected, the access delay

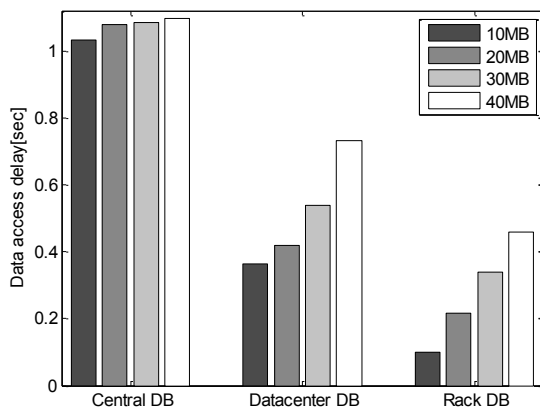


Fig. 6. Data access delay.

becomes smaller for replicas located closer to the servers and for all the replication scenarios an increase in the size of data objects increases in data access delay.

Summarizing, the simulation results presented above show that for cloud applications performing database updates rarely the replication at a closer to computing servers location allows data centers to save energy, conserve bandwidth, and minimize communication delays speeding up execution.

V. CONCLUSIONS AND FUTURE WORK

This paper reviews the topic of data replication in geographically distributed cloud computing data centers and proposes a novel replication solution which in addition to traditional performance metrics, such as availability of network bandwidth, optimizes energy efficiency of the system. Moreover, the optimization of communication delays leads to improvements in quality of user experience of cloud applications.

The performance evaluation is carried out using GreenCloud – the simulator focusing on energy efficiency and communication processes in cloud computing data centers [15]. The obtained results confirm that replicating data closer to data consumers, i.e., cloud applications, can reduce energy consumption, bandwidth usage, and communication delays significantly.

Future work on the topic will be focused developing a formal mathematical model as well as a testbed implementation of the proposed solution.

ACKNOWLEDGMENT

The authors would like to acknowledge the funding from National Research Fund, Luxembourg in the framework of ECO-CLOUD project (C12/IS/3977641).

REFERENCES

- [1] J. G. Koomey, "Worldwide electricity used in data centers," *Environmental Research Letters*, vol. 3, no. 034008, September 2008.
- [2] J. G. Koomey, "Growth in Data center electricity uses 2005 to 2010," Oakland, CA: Analytics Press, August 2011.
- [3] Ruay-Shiung Chang, Hui-Ping Chang, and Yun-Ting Wang, "A dynamic weighted data replication strategy in data grids," *IEEE/ACS International Conference on Computer Systems and Applications (AICCSA)* pp. 414-421, March 2008.
- [4] R. Brown, et al. "Report to congress on server and data center energy efficiency: public law 109-431," Lawrence Berkeley National Laboratory, Berkeley, 2008.
- [5] Li Shang, Li-Shiuan Peh, and N. K. Jha, "Dynamic voltage scaling with links for power optimization of interconnection networks," *International Symposium on High-Performance Computer Architecture (HPCA)*, pp. 91-102, Feb. 2003.
- [6] D. Kliazovich, S. T. Arzo, F. Granelli, P. Bouvry, and S. U. Khan, "Accounting for Load Variation in Energy-Efficient Data Centers," *IEEE International Conference on Communications (ICC)*, Budapest, Hungary, 2013.
- [7] Shengquan Wang, Jun Liu, Jian-Jia Chen, and Xue Liu, "PowerSleep: A Smart Power-Saving Scheme With Sleep for Servers Under Response Time Constraint," *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, vol. 1, no. 3, pp. 289-298, Sept. 2011.
- [8] T. Horvath, T. Abdelzaher, K. Skadron, and X. Liu, "Dynamic voltage scaling in multitier web servers with end-to-end delay control," *IEEE Transactions on Computers*, vol. 56, no. 4, pp. 444-458, 2007.
- [9] M. Guzek, D. Kliazovich, and P. Bouvry, "A Holistic Model for Resource Representation in Virtualized Cloud Computing Data Centers," *IEEE International Conference on Cloud Computing Technology and Science (CloudCom)*, Bristol, UK, 2013.
- [10] D. Kliazovich, J. E. Pecero, A. Tchernykh, P. Bouvry, S. U. Khan, and A. Y. Zomaya, "CA-DAG: Communication-Aware Directed Acyclic Graphs for Modeling Cloud Computing Applications," *IEEE International Conference on Cloud Computing (CLOUD)*, Santa Clara, CA, USA, 2013.
- [11] Bin Lin, Shanshan Li, Xiangke Liao, Qingbo Wu, and Shazhou Yang, "eStor: Energy efficient and resilient data center storage," *International Conference on Cloud and Service Computing (CSC)*, 2011.
- [12] Xiaowen 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, June 2011.
- [13] Fan Ping, Xiaohu Li, C. McConnell, R. Vabbalareddy, and Jeong-Hyon Hwang, "Towards Optimal Data Replication Across Data Centers," *International Conference on Distributed Computing Systems Workshops (ICDCSW)*, pp. 66-71, June 2011.
- [14] Wenhao Li, Yun Yang, and Dong Yuan, "A Novel Cost-Effective Dynamic Data Replication Strategy for Reliability in Cloud Data Centres," *IEEE International Conference on Dependable, Autonomic and Secure Computing (DASC)*, pp. 496-502, Dec. 2011.
- [15] D. Kliazovich, P. Bouvry, and S. U. Khan, "GreenCloud: A Packet-level Simulator of Energy-aware Cloud Computing Data Centers," *Journal of Supercomputing*, vol. 62, no. 3, pp. 1263-1283, 2012.
- [16] D. Chernicoff, "The shortcut guide to data center energy efficiency," Realtimepublishers.com, New York, 2009.
- [17] The Network Simulator Ns2, available at <http://www.isi.edu/nsnam/ns/>
- [18] L. A. Adamic and B. A. Huberman, "Zipf's law and the Internet," *Glottometrics* vol. 3, pp. 143-150, 2002.
- [19] J. Moy, "OSPF Version 2," *IETF RFC 2328*, Apr, 1998.
- [20] X. Fan, W.-D. Weber, and L. A. Barroso, "Power Provisioning for a Warehouse-sized Computer," *ACM International Symposium on Computer Architecture*, pp. 13-23, San Diego, CA, June 2007.
- [21] D. Kliazovich, S. T. Arzo, F. Granelli, P. Bouvry, and S. U. Khan, "e-STAB: Energy-Efficient Scheduling for Cloud Computing Applications with Traffic Load Balancing," *IEEE International Conference on Green Computing and Communications (GreenCom)*, Beijing, China, 2013.
- [22] V. Sivaraman, A. Vishwanath, Z. Zhao and C. Russell, "Profiling per-packet and per-byte energy consumption in the NetFPGA Gigabit router," *IEEE INFOCOM Workshops*, pp.331-336, 2011.
- [23] P. Reviriego, V. Sivaraman, Z. Zhao, J. A. Maestro, A. Vishwanath, A. Sanchez-Macian, and C. Russell, "An energy consumption model for Energy Efficient Ethernet switches," *International Conference on High Performance Computing and Simulation (HPCS)*, 2012.
- [24] Intel Inc. (2010) Intel® Xeon® Processor 5000 Sequence, available at: http://www.intel.com/p/en_US/products/server/processor/xeon5000
- [25] P. Mahadevan, P. Sharma, S. Banerjee, and P. Ranganathan, "Energy aware network operations," *IEEE INFOCOM workshops*, 2009.
- [26] D. Kliazovich, P. Bouvry, and Samee U. Khan, "DENS: Data Center Energy-Efficient Network-Aware Scheduling," *Cluster Computing*, vol. 16, no. 1, pp. 65-75, 2013.