

Energieverbrauch von Live-Migrationen in OpenStack-basierten Private-Cloud-Umgebungen

Christian Pape, Ronny Trommer, Sebastian Rieger

Fachbereich Angewandte Informatik
Hochschule Fulda
Marquardtstraße 35
36039 Fulda
[vorname.nachname]@informatik.hs-fulda.de

Abstract: Die Bedeutung der Energiekosten von Rechenzentren spielt für deren Betreiber eine entscheidende Rolle. Daher wurden zahlreiche Lösungen zur Steigerung der Energieeffizienz im Rechenzentrum eingesetzter IT-Ressourcen entwickelt. Entscheidend ist hierbei u.a. die optimale Auslastung und Verteilung von IT-Ressourcen. Durch Live-Migrationen und Virtualisierungstechnologien können dabei virtuelle Ressourcen im laufenden Betrieb energieeffizient an einem Standort oder darüber hinaus verteilt werden. Der vorliegende Beitrag beschreibt daher ein Verfahren für die Messung des Energieverbrauchs von Live-Migrationen in der OpenStack-basierten Private-Cloud-Umgebung der Hochschule Fulda. Auf Basis der Messungen werden Ansätze für eine Optimierung der Energieeffizienz durch eine standortübergreifende Verteilung der durch OpenStack verwendeten Ressourcen aufgezeigt und evaluiert.

1 Einleitung

Virtualisierungsumgebungen sind in den letzten Jahrzehnten zu einem essentiellen Bestandteil der IT-Infrastruktur von Rechenzentren (RZ) geworden. Durch die Abstraktion der darin betriebenen virtuellen Maschinen (VM) von der darunter liegenden physischen Hardware, wird eine Migration der in den VMs betriebenen Diensten und Anwendungen über mehrere Hosts hinweg ermöglicht. Migrationen lassen sich dabei für Benutzer transparent im Hintergrund und ohne jeglichen Ausfall realisieren (Live-Migration). Migrationen von VMs zwischen Hosts können neben der Verbesserung der Fehlertoleranz auch eine gleichmäßige Verteilung der Auslastung [SR14] der physischen IT-Ressourcen unterstützen. Beispielsweise könnten alle VMs eines gering ausgelasteten Hosts migriert werden, um den Host anschließend abzuschalten, und so den Energieverbrauch zu senken. Private-Cloud-Lösungen, wie z.B. das zunehmend an Bedeutung gewinnende OpenStack, die auf den Virtualisierungsumgebungen aufbauen, bieten zusätzlich Verwaltungs- und Automatisierungsschnittstellen für den Betrieb und die einheitliche Bereitstellung von VMs. Dadurch kann ein Controller in OpenStack-Umgebungen eine energieeffiziente Verteilung der virtuellen Compute-, Storage- und Netzwerk-Ressourcen anstreben. Nicht zuletzt durch Hybrid-Cloud-Lösungen werden zunehmend auch Lösungen interessant, die

eine Verlagerung von VMs (inkl. Storage und Netze) an andere Standorte ermöglichen. Durch diese Lösungen lassen sich Ressourcen und deren Energieverbrauch auf externe RZ verlagern. Aus dem Blickwinkel des Energieverbrauchs lohnt sich eine Migration dann, wenn die Migrationskosten geringer sind als die zu erwartenden Energieeinsparungen. Gemäß [BIT14] haben die steigenden Energiekosten bereits zu einem Abwandern einiger deutscher RZ ins benachbarte europäische Ausland geführt. In der Studie ist der Strompreis pro kWh für RZ in Paris durchschnittlich halb so hoch (7 Cent) als in Frankfurt (14 Cent). Dies unterstreicht die Herausforderung der zunehmenden Energiekosten für RZ sowie den Bedarf für Lösungen zur Steigerung der Energieeffizienz in IKT-Infrastrukturen. In den folgenden Abschnitten der Einleitung werden Einflussfaktoren für den Energieverbrauch in OpenStack-Umgebungen sowie darin unterstützte Migrationsverfahren kurz beschrieben. Abschnitt 2 stellt verwandte Arbeiten vor. Die Messung des Energieverbrauchs anhand der OpenStack-Umgebung der Hochschule Fulda beschreibt Abschnitt 3. Abschnitt 4 zieht ein Fazit und gibt einen Ausblick auf weitere Schritte des Projekts.

1.1 Messung des Energieverbrauchs in OpenStack-Umgebungen

In Private-Cloud-Umgebungen wie z.B. OpenStack können Compute-, Storage- und Netz-Komponenten unterschieden werden. Neben diesen essentiellen Komponenten, die die eigentliche Bereitstellung der Cloud-Dienste übernehmen, bestimmt die erforderliche RZ-Infrastruktur (Klimatisierung, USV, usw.) den Energiebedarf [RGLDA14]. Gemäß einer aktuellen Untersuchung [HC14] existieren in der Bundesrepublik Deutschland ca. 51.100 Rechenzentren (RZ). Allein die Anzahl der großen RZ, d.h. RZ mit mehr als 5000 physischen Servern und Verbrauch im Megawattbereich, stieg zwischen den Jahren 2008 und 2012 um ca. 20%. Im Jahr 2014 belief sich der Stromverbrauch der deutschen RZ auf rund 10 Terawattstunden [Hin15]. Dies entspricht einem Anstieg des Strombedarfs um ca. 3% auch wenn IT-Hardware und RZ-Infrastruktur insgesamt effizienter geworden ist. Setzt sich dieser Trend der Energienutzung von Rechenzentren fort, so ist mit einem Bedarf von ca. 12 Terawattstunden im Jahr 2020 zu rechnen [Hin15]. Das Verschieben von VMs mittels Live-Migration hat einen direkt Einfluss auf die Leistungsaufnahme der genannten Komponenten Compute, Storage und Netz. Der Energieverbrauch der Gesamt-RZ-Infrastruktur lässt sich nur mittelbar durch Migrationen beeinflussen. Compute-Ressourcen werden in der Regel durch die verwendeten Server (CPU, RAM, I/O) bzw. deren Auslastung bestimmt. In den Energieverbrauch des Storage gehen Speichersysteme (Arrays, Controller, HDD, SSD etc.) ein. Die Leistungsaufnahme des Netzwerks wird durch die eingesetzten Netz-Komponenten und ggf. Anzahl der aktiven Links bestimmt. Der konkrete Einfluss von OpenStack auf den Energieverbrauch wird durch dessen Scheduling bzw. Verwaltung der Compute-(Nova), Storage-(Cinder, Swift, Glance) und Netz-Ressourcen (Neutron) bestimmt. In Bezug auf die in diesem Paper untersuchten Live-Migrationen ist beispielsweise entscheidend wie effizient der OpenStack Nova Scheduler die verwendeten Compute-Ressourcen zwischen den an der Migration beteiligten Hosts verlagert. Neben Nova können durch die Verlagerung angebundener Netz- und Storage-Ressourcen z.B. Neutron und Cinder Scheduler Einfluss auf den Energieverbrauch

von Migrationen haben. Im Juno Release von OpenStack ist im Telemetrydienst Ceilometer eine Funktion für die Erfassung von IPMI (Intelligent Platform Management Interface) Daten hinzugekommen, wodurch sich der Energieverbrauch der Compute Nodes rudimentär erfassen lässt. Nicht zuletzt aus dem Blickwinkel der Mess- und Regelungstechnik sollte jedoch die Messung des Energieverbrauchs möglichst mit einem externen Messgerät und nicht auf dem zu messenden System selbst durchgeführt werden. Daher bieten sich intelligente Power Distribution Units (PDUs) an, die zunehmend in Racks eingesetzt werden, um die Leistungsaufnahme angeschlossener Verbraucher zu messen und diese an- bzw. abschalten zu können. Mit dem KiloWatt API (KWAPI) Framework [RGLDA14] existiert eine prototypische Implementierung für die Erfassung der durch externe PDUs und IPMI ermittelten Leistungsaufnahme und deren Integration in Ceilometer. Zusätzlich können die Daten so in übergeordnete Management- und Monitoring-Werkzeuge (vgl. Data Center Infrastructure Management (DCIM)) integriert werden. Eine entsprechende Erfassung der Leistungsaufnahme von Racks (z.B. zur Ermittlung des PUE Werts [WDD12]) bildet die Grundlage für die Steigerung der Energieeffizienz eines RZ und wird daher von RZ-Planern häufig bereits als vorgeschriebenes Energiemanagement umgesetzt. Im Folgenden wird entsprechend eine auf externen PDUs basierende Messung der Leistungsaufnahme von Live-Migrationen in einer OpenStack-Umgebung verwendet.

1.2 Live-Migrationen in OpenStack-Umgebungen

OpenStack bietet verschiedene Funktionen für die Migration von VMs und damit verbundenen Ressourcen. Für die Migration der VMs unterstützt OpenStack Nova unterschiedliche Typen [Fou15]. Dabei wird zwischen *Non-live Migrations* und *Live-Migrations* unterschieden. Bei *Non-live Migrations* werden die Instanzen (VMs) gestoppt, der Zustand auf einen anderen Host (Compute Node) kopiert und dort wieder gestartet. Eine für den Benutzer während des Betriebs der VM transparente *Live-Migration*, wie im Abschnitt 1 genannt, erfordert einige Anpassungen an der OpenStack-Umgebung [Fou15]. Neben dem Aktivieren von `VIR_MIGRATE_LIVE` sollte bei Systemen mit mehreren Netzwerkkarten der Parameter `live_migration_uri` so angepasst werden, dass die bei der Migration übertragenen Daten möglichst über eine separate Netzwerkkarte übermittelt werden. Sofern die Compute Nodes nicht alle über ein gemeinsames Shared Storage (z.B. NFS, Gluster) verfügen, und die Instanzen nicht ausschließlich auf Block Volumes (Cinder Volume-Backed Migration) aufsetzen, die von allen Hosts verwendet werden können, ist ein Kopieren der Storage Volumes bzw. Ephemeral Disks der Instanz erforderlich (*block-migrate*). Alternativ zum Migrieren erlangt das gezielte Stoppen von virtuellen Ressourcen (virtuelle Maschinen, Container, Storage, Netze) auf einem Host und erneute Starten auf einem anderen Host in Private-Cloud-Umgebungen eine zunehmende Bedeutung. Cloud-Ressourcen sind in der Regel darauf ausgelegt ein automatisiertes Spawning (z.B. durch Verwendung von standardisierten Images inkl. Skript für deren Anpassung/Orchestrierung, vgl. *cloud-init*) zu erlauben. Spitze dieser Evolution bilden dabei kurzlebige Instanzen z.B. in Form von Containern (vgl. *Docker*). Für ein standortübergreifendes energieeffizientes Scheduling sind unabhängig davon jedoch nach wie vor Live-Migrations z.B. für

Instanzen erforderlich, die langfristig betrieben werden müssen (z.B. da diese persistente Daten halten vgl. Datenbanken, Storage-, Management-VMs etc.).

2 Verwandte Arbeiten

Die Energieeffizienz bzw. deren Optimierung in Rechenzentren (RZ) und speziell in virtuellen Infrastrukturen werden von einigen wissenschaftlichen Publikationen adressiert. Dabei liegt der Fokus meist auf Algorithmen zur energieeffizienten Platzierung von virtuellen Maschinen (VMs) innerhalb von RZ mit dem Ziel physische Server zu minimieren und der Bildung von Wärmenestern vorzubeugen. Die eigentlichen Kosten von Live-Migrationen werden dabei in der Regel vernachlässigt. Beispielsweise adressieren die Publikationen [BB12] und [SH13] die Platzierung von virtuellen Ressourcen in RZ mit dem Ziel die Auslastung der physischen Server zu optimieren und nicht benötigte Server abzuschalten. Ein verteilter Algorithmus zur Platzierung von VMs in großen Cloud-Umgebungen wird in [WLFJ13] vorgestellt. Die Idee ist hierbei, dass jeder physische Host einen Vektor mit den CPU-Auslastungen der anderen Systeme hält und anhand dessen eine Migrationsentscheidung trifft, um selbst einen unteren bzw. oberen Schwellwert bzgl. der Auslastung einzuhalten. In [VT13] wird ebenfalls ein Algorithmus zur energieeffizienten Platzierung von VMs präsentiert. Basis ist hierbei die Abschätzung des Strombedarfs eines Servers in Bezug auf die Charakteristika der darauf laufenden VMs. Mit VMPlanner [FLL⁺13] und VMflow [MKDK11] existieren darüber hinaus zwei Ansätze, um mittels optimierter Platzierung von VMs nicht die Anzahl der benötigten Server sondern vielmehr die Zahl der benötigten Netzkomponenten zu minimieren. Die vorgestellten Algorithmen liefern damit exemplarisch Ansätze für die energieeffiziente Platzierung von VMs, allerdings wird die dabei genutzte Technik der Live-Migration nicht näher betrachtet oder bewertet. Beispielsweise wird nicht untersucht welche Energieersparnis die Verlagerung einer VM erzielen muss, um die Kosten der eingesetzten Live-Migration zu rechtfertigen. Dies führt schließlich zu der Frage der zeitlichen Granularität dieser VM-Bewegungen, d.h. in welcher Frequenz sind Verlagerungen von VMs noch sinnvoll in Bezug auf den Energiebedarf der beteiligten Komponenten und natürlich auch in Bezug auf die Güte der erbrachten Dienste. Mit [Dar14] existiert ein Artikel der sich genauer mit den Migrationskosten von VMs auseinandersetzt. Hier wird neben den verschiedenen Faktoren, welche die Dauer einer Live-Migration bedingen auch der Einfluss der Reihenfolge aufeinanderfolgender Migrationen untersucht. In [VAN08] wird ein Werkzeug namens pMapper vorgestellt, welches die Platzierung von VMs auf Servern auch unter Berücksichtigung der Migrationskosten zum Ziel hat. Die Autoren versuchen diese Kosten zu bewerten und auch der Einfluss der Migration auf die eigentliche Applikation ist Bestandteil der Betrachtung. Fazit ist hier allerdings, dass die Abschätzung des Energieverbrauchs nach einer erfolgten Migration sehr schwierig ist und daher auch sehr ungenau. Der Schwerpunkt der vorliegenden Publikation liegt daher vorrangig in genau dieser Untersuchung der eigentlichen Migrationskosten im Unterschied zum Normalbetrieb von virtuellen Infrastrukturen. Dabei liefert die Ermittlung des Energiebedarfs auch Aufschluss darüber, welche der beteiligten Komponenten innerhalb der Infrastruktur noch ungenutztes Optimierungspotential bieten.

3 Messung des Energieverbrauchs von Live-Migrationen in der OpenStack-Umgebung der Hochschule Fulda

Die Migration von virtuellen Maschinen erfordert zusätzliche Ressourcen. Um zu erkennen wie sich dieser Prozess auf die Leistungsaufnahme der Compute-, Speicher- und Netzwerkkomponenten auswirkt wurde eine Testumgebung eingerichtet. Abbildung 1 zeigt deren schematischen Aufbau. Als Compute-Knoten werden zwei *Dell PowerEdge R620* mit

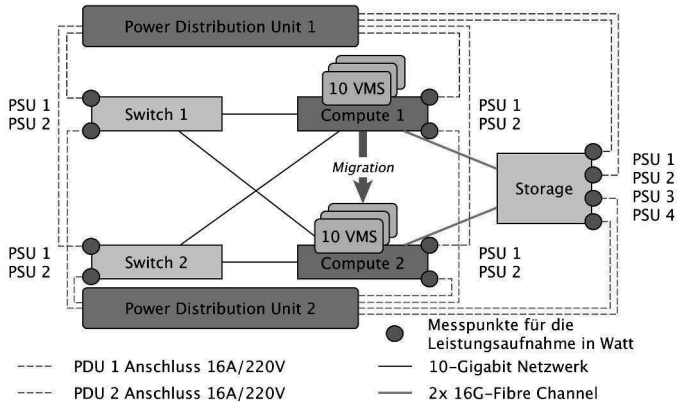


Abbildung 1: Testaufbau und Messpunkte im Versuch

jeweils 2x *Intel Xeon E5-2650* 8-Core CPUs und 256 GB Arbeitsspeicher eingesetzt. Als Netzwerk-Switches wurde ein *Arista 7050S-52* sowie ein *Arista 7150S-24* verwendet. Die Netzanbindung der Compute-Knoten ist jeweils mit 2x 10 Gbit/s gewährleistet. Als Shared Storage wurde ein *NetApp E2700* System verwendet, das an die Compute-Knoten mit je 2x 16 Gbit/s *Fibre Channel* angebunden ist. Um die Leistungsaufnahme unabhängig zu messen wurden zwei IP-fähige *Raritan* Stromverteiler (PDU) vom Typ *PX2-5260R* eingesetzt, deren Messgenauigkeit mit Hilfe eines Digital Power Meter vom Typ *Yokogawa WT333* überprüft wurde. Die festgestellten Abweichungen waren in dem von uns benötigten Messbereich sehr gering und daher vernachlässigbar. Mithilfe der PDUs konnte die Leistungsaufnahme an jedem einzelnen Netzteil gemessen werden.

3.1 Realisierung

Die im vorherigen Abschnitt vorgestellte Testumgebung wurde verwendet um festzustellen welchen Einfluss Migrationen von virtuellen Maschinen (VMs) und damit verbundenen Ressourcen auf die Leistungsaufnahme von Compute-, Speicher- und Netzwerkkomponenten haben. In einem Testverfahren (vgl. Abbildung 2) wurden auf zwei Compute-Knoten VMs eingerichtet, auf einen Compute-Knoten verschoben und anschließend gelöscht. Während des Tests wurden auf den beiden Compute-Knoten keine weite-

1. Leerlauf: Verifizieren, dass keine Ressourcen auf den beiden *Compute-Knoten* anderweitig genutzt werden
2. Messphase für einen Zeitraum von 30 Minuten
3. Initialisierung *Compute-Knoten 1* mit 10 virtuellen Maschinen und ausführen der Simulation für CPU-Last
4. Messphase für einen Zeitraum von 30 Minuten
5. Initialisierung *Compute-Knoten 2* mit 10 virtuellen Maschinen und ausführen der Simulation für CPU-Last
6. Messphase für einen Zeitraum von 30 Minuten
7. Migration der 10 virtuellen Maschinen von *Compute-Knoten 1* auf *Compute-Knoten 2*
8. Messphase für einen Zeitraum von 30 Minuten
9. Löschen der virtuellen Maschinen
10. Leerlauf: Verifizieren, dass keine Ressourcen auf den beiden *Compute-Knoten* anderweitig genutzt werden

Abbildung 2: Testablauf

ren VMs ausgeführt. In den VMs wurde mit der Anwendung *stress* künstliche CPU- und I/O-Last erzeugt. Bei der Einrichtung der Konfiguration der künstlichen Last wurden die Parameter in mehreren Messungen sukzessive angepasst, um ein komplettes Auslasten der Compute-Knoten (*Capping*) zu verhindern. Die Messung der Leistungsaufnahme wird damit nicht durch physische Grenzen limitiert. Zu diesem Zweck wurde die CPU-Auslastung sowie die Linux Kernel-Metrik *Load* über den gesamten Versuchszeitraum gemessen und ausgewertet. Aufgrund der nicht bekannten Leistungskurve der von uns eingesetzten Serverhardware wurde wie in [Sta12] eine energieeffiziente Auslastung von ca. 70%-80% angenommen und mittels *stress* simuliert. Auf die exakte Ermittlung der Leistungskurve mittels Software wurde verzichtet, da der Fokus unserer Untersuchung weniger auf der optimalen Auslastung der physischen Hardware sondern vielmehr auf der Messung der Migrationskosten lag. In dem Versuchszeitraum wurden 10 VMs mit einem *Ubuntu Cloud Image* gestartet. Die VMs erhielten Ressourcen der Kategorie *m1.xsmall* (1 VCPU, 1 GB RAM, 10 GB HDD). Um den Versuch reproduzieren zu können, wurde der Ablauf in Shell-Skripten gegen die *OpenStack Nova API* realisiert. Die Messung der Leistungsaufnahme erfolgte an einer *PDU* direkt an den Netzteilen der Compute-, Speicher- und Netzwerkkomponenten. In Abbildung 1 sind alle Messpunkte der Leistungsaufnahme dargestellt. Alle redundant ausgelegten Netzteile der Komponenten wurden gemessen und in die Auswertung einbezogen. Der Versuchsablauf wurde 3-mal durchgeführt und die Messwerte für die CPU-Auslastung, Systemlast (*Load*) und Leistungsaufnahme verglichen. Die Messdaten über die Systemlast und Leistungsaufnahme wurden mittels *SNMP* und dem freien Netzwerkmanagement-Tool *OpenNMS* in 30 Sekunden Intervallen erfasst und mit dem integrierten *RRDTool* gespeichert, ausgewertet und visualisiert.

3.2 Evaluation

Der zeitliche Verlauf über die Auslastung der beiden Compute-Knoten wird in den beiden Abbildungen 3a sowie 3b dargestellt. In den beiden Abbildungen wurden die unterschiedlichen Zeitabschnitte markiert. Zu den Zeitpunkten A_1 sowie A_2 wurden 10 virtuelle Maschinen (VMs) auf den Compute-Knoten 1 und 2 gestartet. Jede Phase wurde über einen Zeitraum von 30 Minuten gemessen. Um einen besseren Vergleich zwischen der Initialisierung der VMs und der Migration zu ermöglichen, wurden die VMs auf Compute-Knoten 2 um 30 Minuten zeitversetzt gestartet. In den Abbildungen 3a und 3b wird deutlich, dass

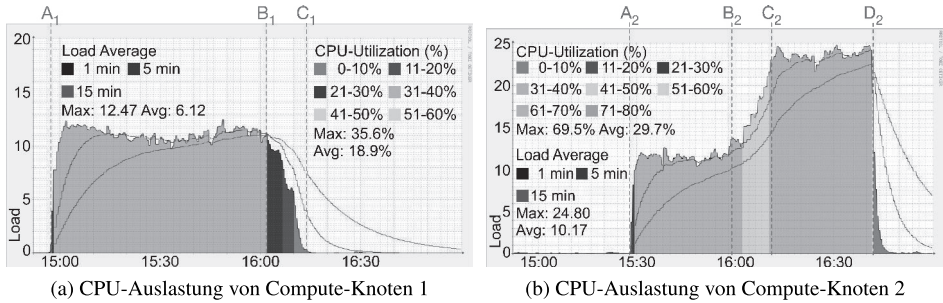
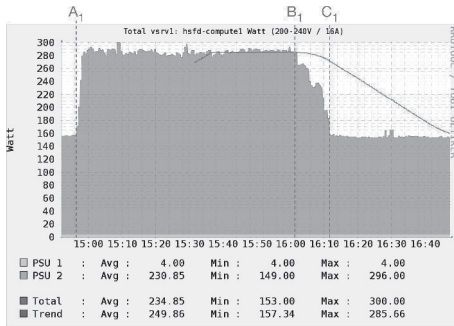
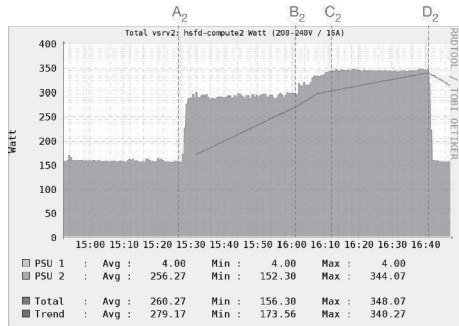


Abbildung 3: Die CPU-Auslastung der beiden genutzten Compute-Knoten

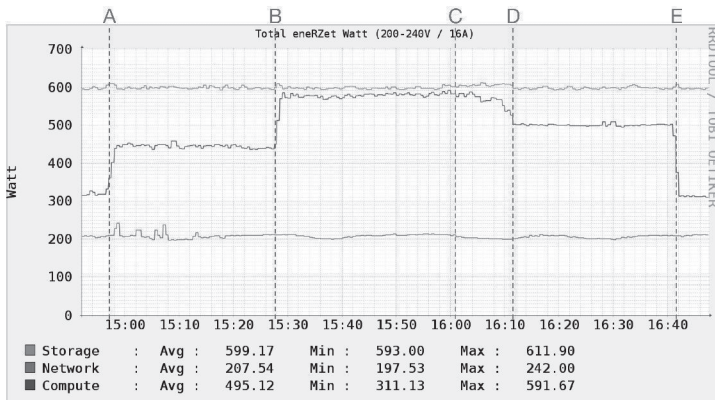
die *Load* zwischen 9.5 und 13 pendelt. Das Lastverhalten mit jeweils 10 VMs ist auf beiden Servern identisch. In den Abbildungen 3a und 3b ist im Abschnitt A der Leerlauf der beiden *Compute-Knoten* sichtbar. Die Gesamtauslastung (*Load*) liegt unter 0,5 und die CPU-Auslastung liegt im einstelligen Prozentbereich. Im nächsten Schritt werden 10 VMs mit simulierter Last auf *Compute-Knoten 1* erzeugt. Wie in der Abbildung ersichtlich ist, wurden zwischen 30% und 40% der verfügbaren CPU-Kapazität genutzt. Die Systemlast steigt auf einen Wert zwischen 10,5 und 12,5. Die Leistungsaufnahme der beiden *Compute-Knoten* verhält sich wie erwartet und steigt auf beiden *Compute-Knoten* von ca. 155 Watt auf ca. 290 Watt an (vgl. Abbildung 4a). Zum Zeitpunkt B_1 (bzw. B_2) wurden 10 VMs von *Compute-Knoten 1* auf *Compute-Knoten 2* verschoben. Die Migration dauert 13 Minuten und die CPU-Auslastung erhöht sich auf dem *Compute-Knoten 2* zwischen 61% und 70%. Die Dauer ergibt sich hierbei zum Einen durch die Migration des Storage der VMs. Zum Anderen wird die Zeit durch Abhängigkeiten in den Migrationsprozessen von OpenStack vorgegeben, die keine gleichzeitige bzw. parallele Migration aller VMs erlauben. Die Systemlast liegt mit 20 VMs bei einer *Load* von ca. 24. Die Anpassung durch Simulation mit dem Kommando *stress* erlaubt eine zuverlässig und reproduzierbare Last um 70%. Mit leicht höheren Parametern wurde durch die Vervielfachung in den virtuellen Maschinen eine höhere Last zwischen 90% und 100% festgestellt. Der Testlauf war mit 70% am besten reproduzierbar, und lag somit zusätzlich nah an der genannten energieeffizienten Auslastung von 80%. Die Leistungsaufnahme von *Compute-Knoten 2* steigt von 280 - 300 Watt auf 350 Watt an (vgl. Abbildung 4b). Nach 30 Minuten Laufzeit wurden die VMs zum Zeitpunkt D_2 in Abbildung 3b gelöscht. Über den Test-Zeitraum wurde zusätzlich die Leistungsaufnahme der Netzwerk- und Storage-Komponenten gemessen. In Abbildung 4c werden alle Messwerte zusammen dargestellt. Die Leistungsaufnahme der Storage-Komponenten (orange) wurde mit durchschnittlich 599 Watt und die Netzwerkkomponenten (grün) durchschnittlich 207 Watt gemessen. Ein interessanter Aspekt ist der Vergleich der Leistungsaufnahme beim Betrieb von 20 VMs mit unterschiedlicher Aufteilung. Werden jeweils 10 VMs auf *Compute-Knoten 1* und *Compute-Knoten 2* verteilt, liegt die Leistungsaufnahme der beiden *Compute-Knoten* zusammen bei ca. 580 Watt. Werden hingegen 20 VMs auf *Compute-Knoten 2* gestartet und *Compute-Knoten 1* läuft im Leerlauf mit, liegt die Leistungsaufnahme lediglich bei ca. 500 Watt. Bei der Initialisierung von 10 VMs stieg die Leistungsaufnahme auf jedem *Compute-Knoten* um ca. 150 Watt.



(a) Leistungsaufnahme Compute-Knoten 1



(b) Leistungsaufnahme Compute-Knoten 2



(c) Gesamte Leistungsaufnahme Server, Netz und Storage

Abbildung 4: Leistungsaufnahme

Jedoch nahm bei der Migration der 10 VMs nach Compute-Knoten 2 die Leistungsaufnahme lediglich um 50 - 70 Watt zu. Im Netzwerk- und Storage-Bereich ist auffällig, dass keine Änderungen der Leistungsaufnahme zu den Zeitpunkten der Initialisierung, Migration und Löschung der VMs zu erkennen sind.

4 Fazit und Ausblick

Live-Migrationen bieten neben Hochverfügbarkeitszenarien zusätzlich Möglichkeiten zur energieeffizienten Nutzung von IT-Ressourcen auch über Rechenzentrumsgrenzen hinweg. In Bezug auf erneuerbare Energiequellen ergibt sich darüber hinaus die Möglichkeit, die Verbraucher in Form von virtuellen Maschinen (VMs) näher an die Stromerzeuger zu bringen. Diese Verlagerung hätte neben ökologischen auch wirtschaftliche Vorteile, wenn durch die Migration günstigerer Strom bezogen werden kann. Die durchgeführten Messungen haben in Bezug auf die von uns eingesetzten Komponenten aufgezeigt, dass

sowohl Storage- als auch Netzwerkkomponenten keinerlei Veränderung im Strombedarf während der Migrationen zeigten, obwohl hierbei sowohl die virtuelle Festplatte als auch der RAM-Inhalt vom Quell- zum Zielsystem übertragen wurden. Daher müssen in unserer OpenStack-Testumgebung sowohl Storage- als auch Netzkomponenten als nicht energieproportional eingestuft werden. Hier bietet sich noch viel Potential für weitere Untersuchungen und Optimierungen. In diesem Zuge wurden bereits die beiden verwendeten Arista Switches hinsichtlich möglicher Erweiterungen untersucht. Durch Anpassung des Linux-basierten Betriebssystems der Switches und Cross-Compilation konnte eine modifizierte Firmware sowie ein neues BIOS (basierend auf coreboot) für die Switches übersetzt und mit Power-Management-Funktionen (ACPI, ASPM) ausgestattet werden. Durch die Erweiterungen lassen sich in dem für die Testumgebung realisierten Layer 2 Multipathing (MLAG) Setup temporär nicht ausgelastete Links abschalten (ca. 2 W/Port mit 10GBase-SR SFP+). Weitere ACPI-Funktionen sind Gegenstand zukünftiger Experimente. Über die Nutzung der EOS CLI per OpenFlow, ist ein adaptives Network Power Management in SDN-Controllern denkbar. Eine Anbindung an das OpenSource Netzwerkmanagementsystem OpenNMS, an dessen Weiterentwicklung zwei der Autoren beteiligt sind, wird hierbei zusätzlich evaluiert. Die präsentierten Ergebnisse bestätigen, dass der Energieverbrauch eines Serversystems nicht linear mit dessen Auslastung ansteigt (vgl. Abbildung 4c). Sofern der Server nicht maximal ausgelastet ist, lassen sich VMs auf einem einzelnen Hostsystem effizienter betreiben, als auf mehreren Hostsystemen. Die dabei potentiell veränderte Dienstgüte (abgesehen von der verringerten Fehlertoleranz) ist Gegenstand unserer nächsten Untersuchungen. Dabei sollen z.B. Webserver-Benchmarks sowohl zur Generierung von Last als auch für die Bewertung der Dienstgüte der VMs zum Einsatz kommen. Es werden auch Situationen betrachtet, in denen eine Migration ineffizient ist und z.B. aufgrund der bestehenden Auslastung einzelner Hostsysteme (vgl. zu wenig verfügbare Cores für zu startende VM) zu weiteren Migrationen führt. Durch die dafür erforderliche Verwendung weiterer Hostsysteme der Private-Cloud der HS-Fulda wird auch die Übertragbarkeit auf reale Rechenzentren mit einer Vielzahl von Hostsystemen und Abhängigkeiten (vgl. Auswirkungen auf Klimatisierung etc.) bewertet. Ein weiterer Ansatz besteht darin anstelle von Live-Migrationen das Scheduling von OpenStack so anzupassen, dass virtuelle Ressourcen energieeffizient verteilt werden. Dabei könnten im Vergleich zu VMs z.B. leichtgewichtige Container (vgl. Docker) unter Berücksichtigung von Storage- und Netz-Abhängigkeiten (z.B. Latenz) am jeweils günstigsten Standort gestartet, und so die Übertragung des Speicherinhalts reduziert werden. Die vorgestellten Live-Migrationen würden dann zum Einsatz kommen, um Instanzen mit persistenten Daten (Datenbank, Management-VM etc.) zwischen Standorten verlagern zu können.

Literatur

- [BB12] A. Beloglazov und R. Buyya. OpenStack neat: A framework for dynamic consolidation of virtual machines in OpenStack clouds. Bericht, CLOUDS-TR-2012-4, Cloud Computing and Distributed Systems Laboratory, The University of Melbourne, 2012.
- [BIT14] BITKOM. Presseinfo EEG und Rechenzentren. http://www.bitkom.org/de/presse/81149_79316.aspx, 2014. 24.4.2015.

- [Dar14] W. Dargie. Estimation of the cost of VM migration. In *Computer Communication and Networks (ICCCN), 2014 23rd International Conference on*, Seiten 1–8. IEEE, 2014.
- [FLL⁺13] W. Fang, X. Liang, S. Li, L. Chiaraviglio und N. Xiong. VMPlanner: Optimizing virtual machine placement and traffic flow routing to reduce network power costs in cloud data centers. *Computer Networks*, 57(1):179–196, 2013.
- [Fou15] OpenStack Foundation. OpenStack Cloud Administrator Guide - Configure migrations. http://docs.openstack.org/admin-guide-cloud/content/section_configuring-compute-migrations.html, 2015. 24.4.2015.
- [HC14] R. Hintemann und J. Clausen. Rechenzentren in Deutschland: Eine Studie zur Darstellung der wirtschaftlichen Bedeutung und der Wettbewerbssituation. http://www.bitkom.org/de/themen/64948_79090.aspx, 2014. 24.4.2015.
- [Hin15] R. Hintemann. Kurzstudie zur Entwicklung von Rechenzentren im Jahr 2014. http://www.borderstep.de/wp-content/uploads/2015/01/Hintemann_Kurzstudie_Rechenzentren_2014.pdf, 2015. 24.4.2015.
- [MKDK11] V. Mann, A. Kumar, P. Dutta und S. Kalyanaraman. VMFlow: leveraging VM mobility to reduce network power costs in data centers. In *NETWORKING 2011*, Seiten 198–211. Springer, 2011.
- [RGLDA14] F. Rossignaux, J.-P. Gelas, L. Lefevre und M. D. De Assuncao. A Generic and Extensible Framework for Monitoring Energy Consumption of OpenStack Clouds. *arXiv preprint arXiv:1408.6328*, 2014.
- [SH13] Nongmaithem Ajith Singh und M Hemalatha. Reduce Energy Consumption through Virtual Machine Placement in Cloud Data Centre. In *Mining Intelligence and Knowledge Exploration*, Seiten 466–474. Springer, 2013.
- [SR14] K. Spindler und S. Rieger. AEQUO - Adaptive und energieeffiziente Verteilung von virtuellen Maschinen in OpenStack-Umgebungen. In *7. DFN-Forum - Kommunikationstechnologien, 16.-17. Juni 2014, Fulda, Germany*, Seiten 45–54, 2014.
- [Sta12] Standard Performance Evaluation Corporation (SPEC). Power and Performance Benchmark Methodology. http://www.spec.org/power/docs/SPEC-Power_and_Performance_Methodology.pdf, 2012. 24.4.2015.
- [VAN08] A. Verma, P. Ahuja und A. Neogi. pMapper: power and migration cost aware application placement in virtualized systems. In *Middleware '08: Proc. of the 9th ACM/I-FIP/USENIX International Conference on Middleware*. Springer, Dezember 2008.
- [VT13] D. Versick und D. Tavangarian. The CÆSARA architecture for power and thermal-aware placement of virtual machines. In *Green Computing Conference (IGCC), 2013 International*, Seiten 1–6. IEEE, 2013.
- [WDD12] M. Wilkens, G. Drenkelfort und L. Dittmar. Bewertung von Kennzahlen und Kennzahlensystemen zur Beschreibung der Energieeffizienz von Rechenzentren. <http://opus4.kobv.de/opus4-tuberlin/frontdoor/index/index/docId/3363>, 2012. 24.4.2015.
- [WLFJ13] X. Wang, X. Liu, L. Fan und X. Jia. A Decentralized Virtual Machine Migration Approach of Data Centers for Cloud Computing. *Mathematical Problems in Engineering*, 2013, August 2013.