

Gesellschaft für Informatik e.V. (GI)

publishes this series in order to make available to a broad public recent findings in informatics (i.e. computer science and information systems), to document conferences that are organized in co-operation with GI and to publish the annual GI Award dissertation.

Broken down into

- seminars
- proceedings
- dissertations
- thematics

current topics are dealt with from the vantage point of research and development, teaching and further training in theory and practice. The Editorial Committee uses an intensive review process in order to ensure high quality contributions.

The volumes are published in German or English.

Information: <http://www.gi.de/service/publikationen/lni/>

ISSN 1617-5468

ISBN 978-3-88579-281-9

The 4. DFN-Forum Communication Technologies 2011 is taking place in Bonn, Germany, from June 20th to June 21st.

This volume contains 12 papers and 4 poster abstracts selected for presentation at the conference.

To assure scientific quality, the selection was based on a strict and anonymous reviewing process.



Paul Müller, Bernhard Neumair, Gabi Dreo Rodosek (Hrsg.): 4. DFN-Forum 2011

GI-Edition

Lecture Notes in Informatics

**Paul Müller, Bernhard Neumair,
Gabi Dreo Rodosek (Hrsg.)**

4. DFN-Forum Kommunikationstechnologien

Beiträge der Fachtagung

20. Juni bis 21. Juni 2011 Bonn

Proceedings





Paul Müller, Bernhard Neumair, Gabi Dreo Rodosek (Hrsg.)

4. DFN-Forum Kommunikationstechnologien

Verteilte Systeme im Wissenschaftsbereich

**20.06. - 21.06.2011
in Bonn**

Gesellschaft für Informatik e.V. (GI)

Lecture Notes in Informatics (LNI) - Proceedings

Series of the Gesellschaft für Informatik (GI)

Volume P-187

ISBN 978-3-88579-281-9

ISSN 1617-5468

Volume Editors

Prof. Dr. Paul Müller

Technische Universität Kaiserslautern

Postfach 3049, 67653 Kaiserslautern

Email: pmueller@informatik.uni-kl.de

Prof. Dr. Bernhard Neumair

Karlsruher Institut für Technologie (KIT)

Steinbuch Centre for Computing (SCC)

Hermann-von-Helmholtz-Platz 1, 76344 Eggenstein-Leopoldshafen

Email: bernhard.neumair@kit.edu

Prof. Dr. Gabi Dreö Rodosek

Universität der Bundeswehr München

Werner-Heisenberg-Weg 39, 85577 Neubiberg

Email: Gabi.Dreö@unibw.de

Series Editorial Board

Heinrich C. Mayr, Alpen-Adria-Universität, Austria (Chairman, mayr@ifit.uni-klu.ac.at)

Hinrich Bonin, Leuphana Universität Lüneburg, Germany

Dieter Fellner, Technische Universität Darmstadt, Germany

Ulrich Flegel, Hochschule Offenburg, Germany

Ulrich Frank, Universität Duisburg-Essen, Germany

Johann-Christoph Freytag, Humboldt-Universität zu Berlin, Germany

Thomas Roth-Berghofer, DFKI, Germany

Michael Goedicke, Universität Duisburg-Essen, Germany

Ralf Hofestädt, Universität Bielefeld, Germany

Michael Koch, Universität der Bundeswehr München, Germany

Axel Lehmann, Universität der Bundeswehr München, Germany

Ernst W. Mayr, Technische Universität München, Germany

Sigrid Schubert, Universität Siegen, Germany

Martin Warnke, Leuphana Universität Lüneburg, Germany

Dissertations

Steffen Hölldobler, Technische Universität Dresden, Germany

Seminars

Reinhard Wilhelm, Universität des Saarlandes, Germany

Thematics

Andreas Oberweis, Karlsruher Institut für Technologie (KIT), Germany

© Gesellschaft für Informatik, Bonn 2011

printed by Köllen Druck+Verlag GmbH, Bonn

Vorwort

Der DFN-Verein ist seit seiner Gründung dafür bekannt, neueste Netztechnologien und innovative netznahe Systeme einzusetzen und damit die Leistungen für seine Mitglieder laufend zu erneuern und zu optimieren. Beispiele dafür sind die aktuelle Plattform des Wissenschaftsnetzes X-WiN und Dienstleistungen für Forschung und Lehre wie die DFN-PKI und DFN-AAI. Um diese Technologien einerseits selbst mit zu gestalten und andererseits frühzeitig die Forschungsergebnisse anderer Wissenschaftler kennenzulernen, veranstaltet der DFN-Verein seit vielen Jahren wissenschaftliche Tagungen zu Netztechnologien. Mit den Zentren für Kommunikation und Informationsverarbeitung in Forschung und Lehre e.V. (ZKI) und der Gesellschaft für Informatik e.V. (GI) gibt es in diesem Bereich eine langjährige und fruchtbare Zusammenarbeit.

Das 4. DFN-Forum Kommunikationstechnologien „Verteilte Systeme im Wissenschaftsbereich“ steht in dieser Tradition. Nach den sehr erfolgreichen Vorgängerveranstaltungen in Kaiserslautern, München und Konstanz wird die diesjährige Tagung vom DFN-Verein und der Universität Bonn gemeinsam mit dem ZKI e.V. und der GI am 20. und 21. Juni 2011 in Bonn veranstaltet und soll eine Plattform zur Darstellung und Diskussion neuer Forschungs- und Entwicklungsergebnisse aus dem Bereich TK/IT darstellen. Das Forum dient dem Erfahrungsaustausch zwischen Wissenschaftlern und Praktikern aus Hochschulen, Großforschungseinrichtungen und Industrie.

Aus den eingereichten Beiträgen konnte ein hochwertiges und aktuelles Programm zusammengestellt werden, das neben künftigen Netztechnologien unter anderem auf virtuelle Forschungsumgebungen, Cloud-Technologien und Identity Management eingeht. Ergänzt wird es durch eine Podiumsdiskussion zu Ergebnissen und Zukunft des Grid-Computing und durch eingeladene Beiträge zu sicherem Cloud-Computing sowie zu Großprojekten im High-Performance-Computing und daraus resultierenden Nutzungsszenarien künftiger Netze. Die Vielfalt und Qualität der für das Forum akzeptierten Beiträge zeigt, dass die Veröffentlichung sowohl im Rahmen der GI-Edition Lecture Notes in Informatics als auch mit Open Access für die Wissenschaftlerinnen und Wissenschaftler attraktiv ist.

Wir möchten uns bei den Autoren für alle eingereichten Beiträge, beim Programmkomitee für die Auswahl der Beiträge und die Zusammenstellung des Programms und bei den Mitarbeiterinnen und Mitarbeitern für die umfangreichen organisatorischen Arbeiten bedanken. Allen Teilnehmern wünschen wir für die Veranstaltung interessante Vorträge und fruchtbare Diskussionen.

Bonn, Juni 2011

Paul Müller
Bernhard Neumair
Gabi Dreo Rodosek

Programmkomitee

Rainer Bockholt, Universität Bonn

Alexander Clemm, Cisco

Gabi Dreo Rodosek (Co-Chair), Universität der Bundeswehr München

Thomas Eickermann, Forschungszentrum Jülich

Markus Fidler, Universität Hannover

Alfred Geiger, T-Systems SfR

Wolfgang Gentzsch, DEISA

Hannes Hartenstein, KIT

Dieter Hogrefe, Universität Göttingen

Eike Jessen, Technische Universität München

Ulrich Lang, Universität zu Köln

Paul Müller (Co-Chair), Technische Universität Kaiserslautern

Bernhard Neumair (Co-Chair), KIT

Gerhard Peter, Hochschule Heilbronn

Christa Radloff, Universität Rostock

Erwin P. Rathgeb, Universität Duisburg-Essen

Helmut Reiser, LRZ München

Peter Schirmbacher, Humboldt-Universität zu Berlin

Uwe Schwiegelshohn, TU Dortmund

Manfred Seedig, Universität Kassel

Marcel Waldvogel, Universität Konstanz

René Wies, BMW Group

Inhaltsverzeichnis

Neue Netztechnologien und Infrastruktur

Raimund Vogl, Norbert Gietz, Markus Speer, Ludger Elkemann

Hochverfügbarkeitsdesign für das Datennetz eines
Universitätsklinikums unter Gesichtspunkten der Anforderungen
klinischer IT-Systeme 15

Andy Georgi, Thomas William, Wolfgang E.Nagel

Synthetische Lasttests auf dem 100-Gigabit-Testbed zwischen der
TU Dresden und der TU Bergakademie Freiberg 25

ITC Management

Sebastian Schmelzer, Dirk von Suchodoletz, Gerhard Schneider

Netzwerkboot über Weitverkehrsnetze: Ansatz zur zentralen
Steuerung und Verwaltung von Betriebssystemen über das Internet 37

Michael Kretzschmar, Mario Golling

Eine Taxonomie und Bewertung von Cloud Computing-Diensten aus
Sicht der Anwendungsentwickler..... 47

Virtuelle Forschungsumgebungen

Martin Gründl, Susanne Naegele-Jackson

Toolset zur Planung und Qualitätssicherung von verteilten virtuellen
Netzwerkstrukturen 59

Andreas Brennecke, Gudrun Oevel, Alexander Strothmann

Vom Studiolo zur virtuellen Forschungsumgebung..... 69

Cloud Computing

Sebastian Rieger, Yang Xiang, Harald Richter

Zugang zu Föderationen aus Speicher-Clouds mit Hilfe von
Shibboleth und WebDAV 81

Pascal Gienger, Marcel Waldvogel

Polybius: Secure Web Single-Sign-On for Legacy Applications 91

Identity Management

Silvia Knittl, Wolfgang Hommel

Identitätsmanagement für Hybrid-Cloud-Umgebungen an
Hochschulen–Erfahrungen im Münchner Wissenschaftsnetz..... 103

Holger Kuehner, Thorsten Hoellrigl, Hannes Hartenstein

Benutzerzentrierung und Föderation: Wie kann ein Benutzer die
Kontrolle über seine Identitätsdaten bewahren? 113

Sebastian Labitzke, Jochen Dinger, Hannes Hartenstein

How I and others can link my various social network profiles as a
basis to reveal my virtual appearance 123

Stephan Wagner, Marc Göcks

IDMS und Datenschutz Juristische Hürden und deren Überwindung . 133

Kurzbeschreibung der Poster

Roland Lemke, Giorgio Flippi, Rolf Chini

EVALSO, an enabling communication infrastructure for Astronomy. 145

Andreas Hanemann, Wolfgang Fritz, Mark Yampolskiy

Dedizierte Wellenlängenverbindungen als reguläres Dienstangebot
der europäischen Forschungsnetze 147

Sebastian Meiling, Dominik Charousset, Thomas C.Schmidt,

Matthias Wählich

HVMcast: Entwicklung und Evaluierung einer Architektur zur universellen
Gruppenkommunikation im Internet.....149

Björn Stelte

Drahtlose Sensornetze in der Lehre-Sicherheitsbetrachtung von

ZigBee-Netzen151

Neue Netztechnologien und Infrastruktur

Hochverfügbarkeitsdesign für das Datennetz eines Universitätsklinikums unter Gesichtspunkten der Anforderungen klinischer IT-Systeme

Dr. Raimund Vogl, Norbert Gietz, Markus Speer, Ludger Elkmann

Zentrum für Informationsverarbeitung
Westfälische Wilhelms-Universität Münster
Röntgenstraße 7 - 13
48149 Münster

rvogl@uni-muenster.de, gietz@uni-muenster.de,
speer@uni-muenster.de, elkeml@uni-muenster.de

Abstract: Das Zentrum für Informationsverarbeitung (ZIV) der Westfälischen Wilhelms-Universität Münster (WWU) betreibt das Kommunikationssystem (Datennetz und Telekommunikation) sowohl der WWU wie auch des Universitätsklinikums Münster (UKM). Dabei wird auf einheitliche technische Standards und Konzepte für die insgesamt ca. 50.000 Netzanschlüsse, von denen ca. 20.000 im Bereich des UKM liegen, gesetzt. Obwohl schon der Netzbetrieb für Forschung und Lehre hohe Anforderungen an die Verfügbarkeit stellt, die auch beim Netzdesign berücksichtigt wurden, stellt der Betrieb in der Krankenversorgung, insbesondere mit zunehmender Verbreitung der elektronischen Krankengeschichte, noch wesentlich weitergehende Anforderungen. Auf Basis betrieblicher Erfahrungen, die speziell die Notwendigkeit für die lokale Eindämmung von Netzstörungen sehr deutlich machen, wurde für die Steigerung der Datennetzverfügbarkeit speziell im Kontext der Einführung der mobilen Visite ein Konzept für die resiliente Layer-3 Entkopplung der DataCenter vom LAN-Core bei gleichzeitig direkter DataCenter-Anbindung einer gesonderten Netzverteilinfrastuktur, die eine besonders robuste Anbindung für selektierte, speziell konfigurierte klinische Arbeitsplätze bieten soll, geschaffen.

1 Die Struktur des Datennetzwerks der WWU und des UKM

Als im Jahr 2000 die Medizinischen Einrichtungen der WWU Münster als UKM zu einer eigenständigen Anstalt wurden, war dies bereits umfänglich mit LAN erschlossen und netztechnisch eng mit dem Netz der Universität verbunden. Insbesondere zentrale Netzfunktionen, Internet-Anschluss, Dokumentation und Betrieb wurden gemeinsam durch das ZIV bereitgestellt. Aus diesem Grunde blieb die Verantwortung für den Netzbetrieb auch des UKM beim ZIV als Dienstleister. Die Konzepte und Standards für das Datennetz, die im Rahmen eines DFG-LAN-Ausbauantrages in 2010 für die Periode 2011-2017 [Ant10, Vog10] neu festgeschrieben wurden, finden auch für das UKM Anwendung.

Die etablierten Eckpunkte des Netzkonzeptes für WWU und UKM können wie folgt zusammengefasst werden:

Performance: Die weit im Stadtgebiet verteilten Liegenschaften von WWU und UKM werden mit einem leistungsfähigen Dark-Fiber-Backbone (insgesamt 250km im Stadtgebiet Münster) mit überwiegend 10GE-Technologie verbunden. Aus Redundanzgründen wird eine Doppel-Stern Topologie (siehe Abb.1) jeweils für WWU und UKM verwendet. Diese getrennten Topologien sind über das für die Vernetzung des gesamten Wissenschaftsstandort Münster aufgebaute WNM (Wissenschaftsnetz Münster, bindet auch FH, MPI, Kunstakademie und das Studentenwerk ein) auf Layer-3 verbunden. Die aktiven Komponenten im Netz des UKM sind in die Schichten Core (Layer-3), Distribution und Access (Layer-2) gegliedert (vgl. Abb.1). Die Verbindungen zwischen den 3 Schichten sind weitgehend in 10 GE realisiert (Reste in 1 GE), die Endgeräteanbindung ist historisch gewachsen (1 GE, 100 MBit/s und in Resten noch 10 MBit/s).

Sicherheit: Für die feingranulare Realisierung von Netzsicherheitszonen und Zugangskontrolle sind an WWU und UKM zusammen ca. 1.000 Endnutzer-VLANs in Betrieb. Damit werden insbesondere auch die parallel abzudeckenden Sicherheits- bzw. Zugangsanforderungen für Forschung und Lehre einerseits und Krankenversorgung andererseits adressiert. Die Layer-3-Funktionalitäten werden durch umfangreichen Einsatz des Cisco IOS-Features VRF-lite (ca. 230 virtuelle Routing-Instanzen) realisiert. Auch die Sicherheitsfunktionen (stateful Firewalling, IPS, VPN) werden virtualisiert in den Netzzonen abgebildet (vgl. [Ant10]). Dabei werden die virtualisierten Sicherheitsfunktionen für WWU und UKM zentral an zwei Netzknoten realisiert. Mit Hilfe des darauf aufsetzenden organisatorischen Netzzonen-Konzeptes lassen sich abgestuft Netzzonen mit an die Nutzungssituation genau angepassten Sicherheitseinstellungen definieren – diese werden vom ZIV gemeinsam mit den Nutzern erarbeitet und reichen von komplett abgeschotteten VLANs für medizinische Spezialsysteme bis zu Bereichen für Forschung&Lehre ohne Einschränkungen des Internetzugangs (abgesehen von den obligatorisch zwischengeschalteten IPS Systemen).

Verfügbarkeit: Es wird primär auf Geräte-Dopplung (an zwei getrennten Standorten, mit getrennt geführten Kabelwegen) und nicht auf intrinsische Redundanz in den Geräten gesetzt. Für Redundanzfunktionen auf Layer-2 wird derzeit noch Spanning Tree in verschiedenen Varianten (RSTP und PVST+) sowie auf Layer-3 HSRP und OSPF eingesetzt.

Management und Dokumentation: Ein im Hause entwickeltes, integriertes Datenbank- und Anwendungssystem wird für Konfigurationsverwaltung, Workflow-Automatisierung, Case Management, LAN-Dokumentation und Bauprojektanbahnung (LAN-Base) eingesetzt. Die Netzüberwachung wird mit dem Produkt CA Spectrum realisiert.

Service: Für WWU und UKM gemeinsam werden ein NOC (Network Operating Center für die Netzbetriebsführung), ein NIC (Network Information Center für die

Netzanschluss-, Endgeräte- und Adressverwaltung), ein NTC (Network Technology Center für Lagerverwaltung) sowie eine gemeinsame 365x24 Rufbereitschaft angeboten.

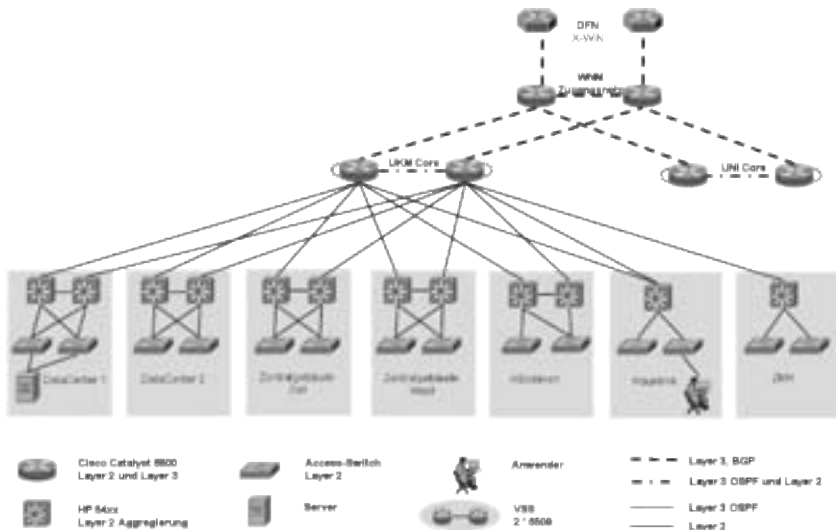


Abb.1: Ist-Situation des Netzes des UKM - Ein Core mit 2 redundanten Switches bildet das Zentrum des Doppelsterns der Layer-2 Verteilstruktur. Auch die beiden DataCenter sind auf Layer-2 direkt an die Core-Switches angebunden. Die Betriebserfahrung zeigt, dass sich dadurch leicht Störungen über den Core in den sensiblen DataCenter-Bereich ausbreiten. Die Verbindung zum Netz der WWU erfolgt über Layer-3 über die Router des Wissenschaftsnetzes Münster (WNM). Hierüber werden auch weitere Hochschulen und Forschungseinrichtungen am Standort eingebunden und die Konnektivität zum DFN hergestellt.

2 Betriebliche Erfahrungen – Bedarf für die lokale Eindämmung von Störungen

Das oben beschriebene Netzdesign wurde seit 2002 aufgebaut und seitdem stetig von anfänglich ca. 23.000 auf gegenwärtig ca. 50.000 Netzanschlüsse im gemeinsamen Netz von WWU und UKM erweitert. In dieser Zeit ist es zu mehreren umfänglicheren Störungen des Netzes gekommen, die im Case Management Tool umfänglich dokumentiert sind.

Ein zentraler Aspekt bei den betrieblichen Erfahrungen der letzten Jahre ist dabei, dass die aus Gründen der hohen Flexibilität bei der Bereitstellung von Netzanschlüssen in derselben Netzzone für auf verschiedene Liegenschaften verteilte Arbeitsgruppen und der feingranular zentral bereitstellbaren Sicherheitsfunktionen etablierte Strategie der geografie-unabhängigen Bereitstellung von Endnutzer-VLANs problematisch ist. Auch die uneingeschränkte Zentralisierung der Netzsicherheits- (stateful Firewall, IPS, VPN)

und Kerndienstfunktionen (DHCP, DNS, ...) steht mit wachsender Netzausdehnung im Konflikt mit den drastisch gestiegenen Anforderungen an die Netzverfügbarkeit.

So ist festzustellen, dass beispielsweise die Absicherung von sensiblen Netzzonen einzelner Bereiche im Klinikum durch zentral realisierte Firewall und IPS Funktionen bei Störungen des WWU Netzes durch hohe Netzbelastung (z.B. DoS Attacken) deren netzmäßige Isolierung bewirken, obwohl im UKM Netz ansonsten keine Störungen bestehen. Ein ähnliches Bild stellt sich auch bei Störungen und Fehlkonfigurationen der zentralen Sicherheits- und Kerndienstfunktionen dar.

Wiederholt war auch zu bemerken, dass Schleifen im Netzwerk (wegen versehentlicher „Kurzschluss-Patchungen“ durch Endnutzer) – entgegen der Erwartung in die zu deren Erkennung und Unterdrückung eingesetzten Funktionalitäten der Netzkomponenten – immer wieder zu großflächigen Störungen führen, die insbesondere innerhalb des gesamten betroffenen Layer-2 Bereichs (also im dargestellten Netz-Design potenziell campusweit) ausgebreitet werden, und durch Überlastung aller davon betroffenen Switches auch Kollateraleffekte in anderen Layer-2 Bereichen haben können. Die Überlastung der switchinternen Controlplane bewirkt ein willkürliches Verwerfen von Steuerpaketen, in dessen Folge verschiedenste Protokolle zur Kontrolle und Redundanzbereitstellung kollabieren (beispielsweise die Routingprotokolle - OSPF, STP) oder gar die Links der Coreswitches wegen des Verwerfens der UDL (Unidirectional Link Detection) Pakete in den Zustand Errordisable schalten.

Entsprechende großflächige Störungen waren Anfang 2010 zu bewältigen, eine niemals aufgeklärte Störung Anfang 2008 dürfte eine ähnliche Ursache gehabt haben. Ebenso können nicht beeinflussbare Störquelle wie ein defektes Switchmodul in einem der Coreswitches in April 2010 durch ein Fehlverhalten mit falsch weitergeleiteten Paketen zu campusweiten Störungen führen. Auch zeitlich sehr beschränkte Störungsquellen führten dabei zu schmerzhaft langen Netzausfällen, bis sämtliche Links wieder aktiviert und die Kontrollprozesse wieder stabilisiert werden konnten. Die Fehlersuche in solchen großen Störungssituationen gestaltet sich nicht nur durch den dann erheblichen Druck auf die Mitarbeiter als äußerst schwierig, auch sind viele Netzkomponenten wegen der Netzprobleme nicht mehr erreichbar und durch die hohe CPU-Auslastung kaum noch bedienbar. Die Bedeutung eines umfassenden und verlässlich verfügbaren out-of-band Management für die Core-Switches wird dabei deutlich. Auch Wartungsmaßnahmen (z.B. Softwareupdates) stellen in der aktuellen betriebenen Netzarchitektur ein großes Störpotenzial dar, führen doch Reboots immer (bedingt durch die Layer-2 Ausdehnung und der Layer-3 Konzentration) zu campusweiten Störungen des IT-Betriebes, die durch ihre Rückwirkungen auf die Anwendungssysteme (z.B. Clusterschwenks und hängende Dienste) deutlich über die kurzzeitige Downtime des LAN hinausgehen.

Ein im Oktober 2010 gemeinsam mit Gerätehersteller Cisco, dem ZIV und dessen Servicepartner durchgeführtes Customer Proof-of-Concept (CPOC) in den Cisco-Labors in London zeigte, dass in dem vorhandenen Netzdesign keine verlässlichen Funktionen der Netzwerkkomponenten genutzt werden können, um auftretende Netzstörungen in den lokalen oder funktionalen Bereich ihrer Entstehung zu begrenzen. Hierbei zeigte sich auch für den Hersteller Cisco unerwartet die generelle Empfindlichkeit von

Netzwerkkomponenten und Protokollen bei einer starken Aus- und Überlastung der Controlplane. Verschiedenste Versuche durch steuernden Eingriff in das Aufkommen von Control-Datenpaketen einzugreifen oder deren Transport auf der Controlplane zu priorisieren brachten keine nennenswerte Verbesserung der Netzstabilität oder führten zu anderen mindestens ebenso schmerzlichen neuen Verwundbarkeiten des Netzes. Einzig eine drastische Reduzierung der Komplexität im bestehenden Netzdesign durch Beschränkung der Anzahl der VLANs auf eine kleine zweistellige Zahl gewährte einen ausreichenden Schutz und zeigte, dass die eingesetzten Verfahren und Protokolle in der Lage sind, die Störungen abzufangen bzw. zu begrenzen.

Offensichtlich besteht also trotz durchdachtem Netzdesign und gut organisierter Betriebsführung eine gefährliche, nicht vermeidbare Verletzlichkeit des Datennetzwerkes als Ganzes durch geringfügige lokale Störungsursachen (z.B. Schleifen im Endnutzerbereich). Auch wenn durch einzelne Konfigurationsmaßnahmen (aktualisierte Softwarestände mit Bugfixes im Bereich Ressourcen-Allozierung auf den Switches; Automatisches Wiederanlaufen nach Errordisable; Priorisierung von Steuerpaketen auf der Controlplane; Prävention von Schleifen durch Deaktivierung der bequemen auto-cross-out Funktion) kurzfristig eine Verbesserung bei der Verwundbarkeit möglich ist, besteht ein Bedarf für eine grundlegende Überarbeitung des Netzdesigns, speziell mit Augenmerk auf die mit wachsender Durchdringung des klinischen Bereiches mit IT-gestützten Verfahren und der integrierten Nutzung des Datennetz für die Kommunikation (z.B. VoIP) immer bedeutsamer werdenden Anforderung für höchste Verfügbarkeit.

3 Das klinische Anforderungsszenario – die mobile Visite

Das UKM plant im Zuge der Etablierung der elektronischen Krankengeschichte die Einführung der mobilen Visite auf sämtlichen Bettenstationen. Dabei soll, gestützt über das dafür flächendeckend zu etablierende WLAN, der Abruf und die Erfassung von Krankengeschichtsdaten (Vitalparameter, Medikation, etc...) über mobile Geräte direkt am Patientenbett realisiert werden. Durch den geplanten Verzicht auf die Papier-Form wird die „100%-ige“ Verfügbarkeit der elektronischen Krankengeschichte unverzichtbar für die adäquate Behandlung.

Wegen der hohen Belastungen und des Erfolgsdruckes in diesem ambitionierten Projekt sieht sich insbesondere auch das ZIV als Betreiber des dafür unverzichtbaren Datennetzes mit der verständlichen Anforderung konfrontiert, alle Möglichkeiten zur Sicherstellung der höchstmöglichen Netzverfügbarkeit (Forderung: 100%) auszuschöpfen. Die obige Darstellung der in den letzten Jahren offenkundig gewordenen Verwundbarkeiten des Datennetzwerkes macht deutlich, dass dafür neue technische Designkonzepte notwendig sind.

Nach der Abwägung zahlreicher auch unorthodoxer Lösungsansätze (z.B. Rückgriff auf externe Mobilfunknetze) zur Schaffung von redundanten Systemen zur Bewältigung von – effektiv unvermeidbaren – Störungen konnte das im nächsten Abschnitt beschriebene Redesign für die Schaffung selektierter hochverfügbarer klinischer Arbeitsplätze

erarbeitet und mit der Nutzerseite akkordiert werden, das folgenden zentralen Forderungen genügt:

1. Das Konzept soll einfach und überschaubar sein. Komplexe und nicht unter eigener Kontrolle stehende Technologien und Systeme (wie z.B. bei einem Backup über externe 3G Netze) sind genauso ungeeignet wie neue, betrieblich noch nicht erprobte Komponenten.
2. Im Störfall muss für selektierte „hochverfügbare“ Arbeitsplätze ein nahtloses Weiterarbeiten ohne die Notwendigkeit von Interventionen zur Umkonfiguration möglich sein.
3. Die hochverfügbaren Arbeitsplätze müssen in täglicher Routineverwendung stehen, da nur so die Funktion im Krisenfall sichergestellt ist.
4. Der Kommunikationsweg zwischen den hochverfügbaren Arbeitsplätzen und den DataCenters, in denen die Anwendungssysteme der elektronischen Krankengeschichte betrieben werden, muss möglichst direkt sein.

4 Hochverfügbarkeits-Redesign für das Datennetz

Auf Basis der betrieblichen Erfahrungen der letzten Jahre hatte sich die großräumige Ausbreitung von Störungen mit dezentral lokalisierbaren Ursachen als Kernproblem herausgestellt. Eine grundlegende Erneuerung des gesamten LAN, die sowieso an der Zeit ist und in deren Zuge auch diese Problematik adressiert werden soll, ist mittelfristig geplant. Kurzfristig ist eine Abkehr vom etablierten VLAN-Konzept im UKM Campus Backbone mit seinen ca. 20.000 Netzanschlüssen jedoch nicht bewerkstellbar.

Die netztechnische Grundlage für die Hochverfügbarkeit der zentralen klinischen Informationssysteme der elektronischen Krankengeschichte an einer überschaubaren Zahl (wenige 100) von selektierten und funktional eingeschränkten klinischen Arbeitsplätzen lässt sich jedoch mit überschaubarem Aufwand (in puncto Material, Zeit und Arbeitsleistung) mit folgenden Maßnahmen realisieren.

Layer-3 Entkopplung der DataCenter vom Campus-Backbone

Die Ausbreitung von Störungen im normalen Campusnetzwerk in die DataCenter muss unterbunden werden. Nur so kann die Verfügbarkeit der dort betriebenen IT-Systeme der elektronischen Krankengeschichte sichergestellt werden. Die Entkopplung auf Layer-3 durch ein eigenes Switch-Paar ist eine Maßnahme dafür, die sich anhand der betrieblichen Erfahrungen bewährt hat – Störungen im WWU Netz wurden durch die Layer-3 Entkopplung im WNM effektiv vom UKM Netz ferngehalten. Der Materialaufwand dafür hält sich in Grenzen. Die damit verbundene Möglichkeit zur Etablierung neuer „virtual switching“ Technologien (VSS: Virtual Switching System), die aktiv/aktiv Redundanz ohne von STP bekannten langen Konvergenzzeiten ermöglicht, und somit die Perspektive zur Eliminierung der Cluster-Schwenks bei den

hochverfügbar geclusterten und auf die beiden DataCenters verteilten Systemplattformen der klinischen IT-Systeme bietet, ist ein höchst willkommener Zusatznutzen. Der Vergleich von Abb.1 und Abb.3 verdeutlicht die Umstellungen im LAN zur Entkopplung der DataCenter. Die in Abb.2 und Abb.3 dargestellten Strukturen wurden ebenfalls im oben beschriebenen Cisco-CPOC geprüft und boten hier den einzigen verlässlichen Schutz vor einer Übertragung von Störungen im vorhandenen Netz auf den DataCenter-Bereich. Es muss hierbei nicht auf die Bildung von Netzzonen und deren Entkoppelung per VRF-lite verzichtet werden. Es besteht die Möglichkeit eines virtualisierten Backbones im Core, über den die einzelnen VRFs mittels Transfer-VLANs verbunden werden. Diese dürfen sich jedoch lediglich als Punkt-zu-Punkt-VLAN auf einzelnen Links befinden, ein Weiterleiten dieser Transfer-VLANs durch Geräte hinweg ist im ersten Schritt nicht mehr zulässig (siehe Abb.2). Die Erfahrungen aus den CPOC zeigen zusätzlich, dass eine Beschränkung auf wenige VLANs in dem entkoppelten Netzbereich die Netzstabilität - auch in extremen Störungssituationen innerhalb der DataCenter - erhöhen kann. Auch wenn keine exakte Menge ermittelt wurde (oder ermittelt werden kann) scheint eine Anzahl von nicht wesentlich mehr als 30 VLANs, die sich in unserem Fall sinnvoll auf 2-3 VRFs-lite verteilen lassen, sinnvoll.

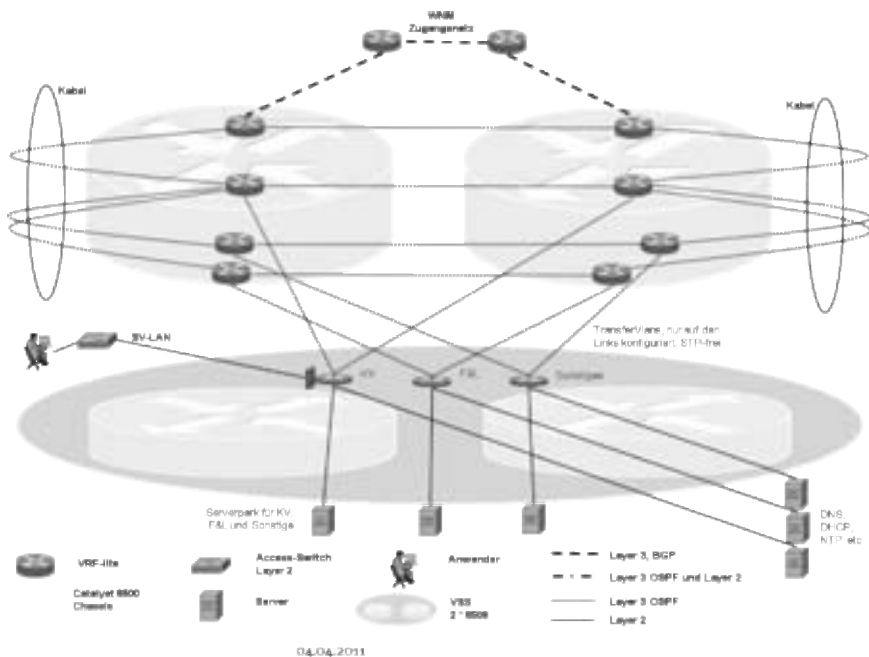


Abb.2: Netzzonen können weiterhin mittels VRF-lite voneinander getrennt werden. Die Verbindung der VRF-lite erfolgt über ein virtualisiertes Backbone mittels link-begrenzten Transfer-VLANs.

Eine wichtige Konsequenz bei der Umsetzung dieser Layer-3 Entkopplung ergibt sich jedoch für den Betrieb von Client-Server Systemen. Da die selben VLANs nicht mehr im Campusnetzwerk und im DataCenter definiert sein können, müssen die

Kommunikationsbeziehungen zwischen Clients und Servern aufwändig neu organisiert werden – generell müssen neue IP Service-Adressen konfiguriert werden, und nicht gerouteter Datenverkehr zwischen Clients und Servern ist nicht mehr möglich.

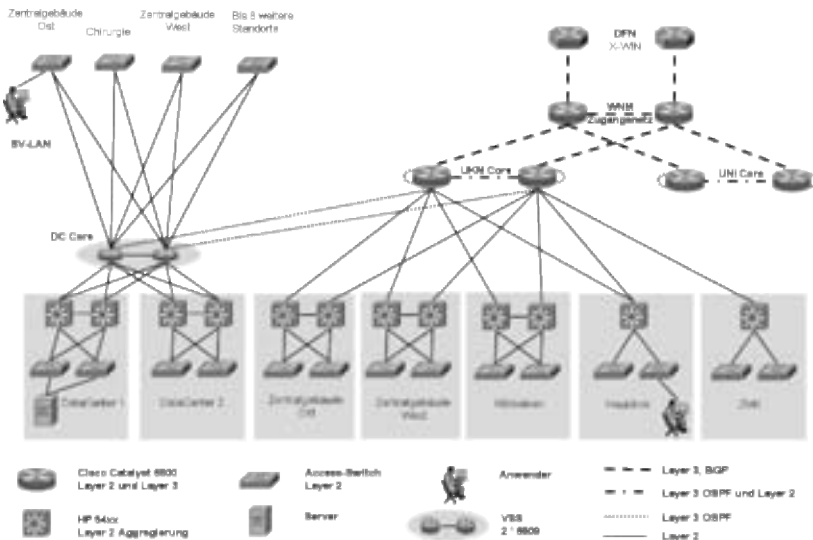


Abb.3: Soll-Konzept für das Netz des UKM - Die DataCenter werden auf Layer-3 vom Core des Campusnetzes entkoppelt, um die Fortpflanzung von Störungen in diesen sensiblen Bereich zu verhindern. Die redundante Anbindung der Verteilstruktur im DataCenter wird anstelle aktiv/passiv (Spanning Tree) über eine aktiv/aktiv Link-Aggregation Verbindung an die über „virtual switching“ (VSS) gekoppelten DataCenter-Switches realisiert. Eine direkt an die DataCenter Switches angebundene hochverfügbare Verteilstruktur (SV-LAN) für selektierte Arbeitsplätze reduzierter Funktionalität stellt auch im Falle großflächiger Netzstörungen im Rest des UKM-Netzes deren Verbindung zu den DataCenters sicher.

Gesondertes Verteilnetz mit enger Anbindung an die DataCenter

Für die Verfügbarkeit der elektronischen Krankengeschichtssysteme in den klinischen Bereichen ist die Datennetzkonnektivität von dortigen Arbeitsplatzsystemen zu den Servern in den DataCenters notwendig. Da Störungen im „normalen“ Campus-LAN des UKM – wie zuvor dargestellt - realistisch nicht ausgeschlossen werden können und eben deswegen die vorstehend beschriebene Abkopplung der DataCenter umzusetzen ist, ist für hochverfügbare Arbeitsplatzsysteme auch ein vom „normalen“ Campus-LAN unabhängiger Kommunikationsweg notwendig. Dieser muss direkt zu den DataCenter-Switches führen und den normalen Campus-Backbone umgehen. Dies bedingt auch direkte passive Netzverbindungen von Access-Switches an die DataCenter-Switches. Hierfür kann glücklicherweise auf eine multimode Fiber-to-the-Desk-Verkabelung aus früheren Jahren, mit denen anfänglich wegen langer Wegstrecken in den Bettentürmen des UKM Arbeitsplätze mit Datennetz versorgt wurden, zurückgegriffen werden. Abseits des seit längerem auf CAT6 Kupferkabeln basierenden normalen Campus-LAN

lassen sich damit wenige selektierte Arbeitsplatzsysteme (ca. 2 bis 3 je Station, ca. 250 insgesamt) auf sehr direktem Weg an wenige Access-Switches an den DataCenters anbinden (vgl. Abb.3). Durch das weitestgehende Vorhandensein der passiven Infrastruktur und den geringen Bedarf an aktiven Komponenten hält sich auch hier der Materialaufwand in engen Grenzen. In Analogie zur Energieversorgung wird dieses LAN-Segment als Sonderversorgungs-LAN (SV-LAN) bezeichnet.

Selektierte funktional eingeschränkte klinische Arbeitsplätze

Die Anbindung der ca. 250 ausgewählten Arbeitsplatzsysteme erfolgt sehr direkt an die DataCenter-Netze. Aus offensichtlichen Gründen der Verfügbarkeit muss hier auf den Rückgriff auf die zentralen virtualisierten Netzsicherheitsfunktionen (stateful FW, IPS, VPN) verzichtet werden – allenfalls können ACLs und Edge-Portbasierende Funktionen (MAC-Security, DHCP-Snooping, etc) genutzt werden. Folglich muss mit diesen Arbeitsplätzen sehr restriktiv bzgl. Sicherheit verfahren werden. Der möglichst ausschließliche Einsatz von ThinClients, die auf die klinischen IT-Systeme nur im Wege der ebenfalls in den DataCenters betriebenen Citrix Terminalserver zugreifen, gewährleistet hohe Sicherheit und optimiert durch die Minimierung von clientseitigen Störungen weiter die Verfügbarkeit. Da diese Arbeitsplätze aber voll für die tägliche Routinearbeit genutzt werden können, ist sichergestellt, dass clientseitige Störungen sofort bemerkt werden und diese Geräte somit im Krisenfall wesentlich verlässlicher bereitstehen als ausschließlich auf standby bereitgehaltene Notlösungen. Darüber hinaus evaluiert der für die Betreuung dieser Arbeitsplatzsysteme zuständige Geschäftsbereich IT des UKM Verfahren zur zentralen Überwachung der Citrix-Clients sowie der notwendigen Serverdienste und Schnittstellen. Die Bereitstellung einer dedizierten Citrix-Farm für die Clients im SV-LAN wird ebenfalls diskutiert.

5 Zusammenfassung und Ausblick

Das dargestellte Konzept verspricht eine wesentliche Erhöhung der Verfügbarkeit der klinischen IT-Systeme der elektronischen Krankengeschichte im UKM durch Verbesserung der Netzkonnektivität für selektierte Arbeitsplatzsysteme in den klinischen Bereichen. Durch die moderaten Umsetzungsaufwände konnte eine schnelle Umsetzung bereits in Angriff genommen werden. Die erhofften Verfügbarkeitsverbesserungen sollen durch Monitoring (sowohl über die Netzmanagement-Systeme wie auch durch Beobachtung und Evaluation gemeinsam mit den Nutzern) verifiziert werden.

Das dargestellte Designprinzip der Layer-3 Entkopplung soll bei der anstehenden Erneuerung des LAN-Backbone von WWU und UKM weitergehend angewandt werden und damit ultimativ auch im „normalen“ Campus-LAN des UKM die Verfügbarkeit erhöhen.

In diesem Konzept ist die deutliche Abkehr von einem in der bisher etablierten Netzarchitektur gültigen Prinzip erkennbar. Wurde bislang jeder Layer-2 Bereiche global im ganzen Netz angebunden, die Layer-3 Vermittlung zwischen diesen aber nur an

wenigen zentralen Core-Routern, sollen zukünftig die Layer-2 Bereiche im Sinne des Störungs-Containments lokal beschränkt werden, und die Vermittlung in andere Layer-2 Bereiche durch dezentrale Routing-Instanzen erfolgen, die ihren Datenverkehr untereinander über wenige Transfer-VLANs über wenige leistungsfähige zentrale Layer-2 Switches austauschen.

Selbstverständlich hat dies auch substantielle Auswirkungen auf die Betreiber von Client-Server-Systemen, da die bislang mögliche Bereitstellung der Endnutzer-Ports auch im DataCenter nicht mehr möglich ist, und auch im Hinblick auf die neuen Konzepte des DataCenter Bridging (DCB) nicht mehr zeitgemäß ist.

Literaturverzeichnis

- [Ant10] Antrag zum Ausbau des Kommunikationssystems – Netzkonzept, Netzentwicklungsplan, Betriebs- und Managementkonzept, Personalsituation – Gemeinsame Darstellung für die Westfälische Wilhelms-Universität und das Universitätsklinikum Münster. Münster, 2010.
http://www.uni-muenster.de/imperia/md/content/ziv/pdf/antrag_ausbau_kommunikationssystem_teil_b-version16-release_mit_anhaengen.pdf.
- [Vog10] Vogl, R; Speer, M; Gietz, N; Elkemann, L.: Netzentwicklungskonzept für ein großes Universitätsnetzwerk – Bestandspflege und Erschließung neuer Technologien. In (Müller, P. et.al. Hrsg.): Proceedings 3. DFN-Forum Kommunikationstechnologie, Konstanz 2010, Gesellschaft für Informatik, Bonn, 2010; S. 45-60.

Synthetische Lasttests auf dem 100-Gigabit-Testbed zwischen der TU Dresden und der TU Bergakademie Freiberg

Andy Georgi, Thomas William, Prof. Dr. Wolfgang E. Nagel

Andy.Georgi@tu-dresden.de
Thomas.William@zih.tu-dresden.de
Wolfgang.Nagel@tu-dresden.de

Abstract: Das Hochleistungsrechenzentrum der TU Dresden wurde in Kooperation mit Alcatel-Lucent und T-Systems über eine 100-Gigabit-Strecke mit dem Rechenzentrum der TU Bergakademie Freiberg verbunden. Die Umsetzung erfolgt erstmals mit kommerzieller Hardware. Eine weitere Besonderheit ist die Übertragung der Daten über eine einzige Wellenlänge. In dieser Arbeit werden nach der Beschreibung des 100-Gigabit-Testbeds erste Ergebnisse synthetischer Lasttests auf verschiedenen Ebenen des OSI-Referenzmodells vorgestellt und diskutiert. Die Vorgaben und der spezielle Aufbau des Testbeds machen eine Überarbeitung vorhandener Messmethoden notwendig. Die Entwicklung der Burst-Tests wird ebenfalls in dieser Arbeit beschrieben. Zusätzlich wird ein Ausblick auf weitere geplante Teilprojekte, sowie Einsatzmöglichkeiten der 100-Gigabit-Technik in Forschung und Industrie gegeben.

1 Einführung

Die verfügbare Bandbreite der Datennetze wird durch aktuelle Trends wie Grid- und Cloud-Computing, aber auch durch die stetig steigende Menge an zu verteilenden Daten vielfach zum bestimmenden Faktor. Für Wissenschaft und Forschung werden daher Gigabit-Netze in Zukunft unerlässlich. Die Erprobung der neuen 100-Gigabit-Technologie zwischen Dresden und Freiberg unter realen Bedingungen ist dafür ein wichtiger Meilenstein.

Alcatel-Lucent und T-Systems haben das Hochleistungsrechenzentrum der TU Dresden und das Rechenzentrum der TU Bergakademie Freiberg im Rahmen eines Testbeds miteinander vernetzt. Es kommt erstmals kommerzielle 100-Gigabit-Technik mit aufeinander abgestimmten IP-Routern und optischen Übertragungssystemen zum Einsatz. Diese Kombination ermöglicht eine besonders hohe Skalierbarkeit, Kapazität und Kompatibilität in Core-, Edge- und Metro-Netzen. Die Daten auf der 100-Gigabit-Strecke zwischen den beiden Technischen Universitäten werden in Echtzeit verarbeitet und über eine einzige Wellenlänge übertragen.

Dank ihrer hohen Bandbreite ermöglicht die 100-Gigabit-Technik den Forschungsinstituten eine neue Qualität der Zusammenarbeit. Sie eignet sich für HPC- und Cloud-basierte

Dienste sowie für Multimedia-Anwendungen, welche das gemeinsame Arbeiten von Nutzern unterstützt (Kollaboration). Die neue Verbindung bietet eine Umgebung für die Evaluierung der Eigenschaften von bandbreitenintensiven Anwendungen und Diensten unter realen Bedingungen. Wissenschaftler der beiden Universitäten untersuchen dabei den Ressourcenbedarf von HPC-, Multimedia- und verteilten Server-Anwendungen.

In dieser Arbeit werden nach der Beschreibung des 100-Gigabit-Testbeds (Kapitel 2) erste Ergebnisse synthetischer Lasttests auf verschiedenen Ebenen des OSI-Referenzmodells vorgestellt und diskutiert (Kapitel 4). In Zusammenarbeit mit den Herstellern wurden konkrete Testszenarien definiert und implementiert, welche in Kapitel 3 erläutert werden. Dabei soll untersucht werden, welche Effekte der Burst-Tests auf TCP/IP Ebene bei Verwendung burst-artiger I/O Muster auf der Anwendungsebene noch sichtbar sind. Zum Abschluss erfolgt in Kapitel 5 eine Zusammenfassung sowie ein Ausblick auf weitere geplante Teilprojekte auf dem 100-Gigabit-Testbed und Einsatzmöglichkeiten in der Forschung und Industrie.

2 Aufbau der Teststellung

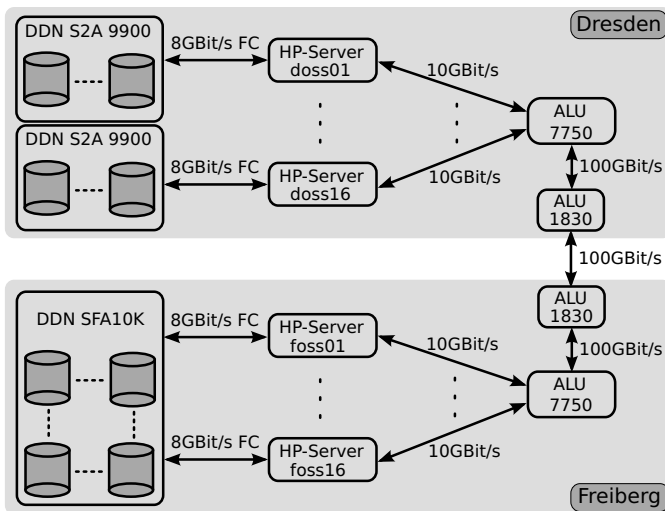


Abbildung 1: Aufbau der Teststrecke

Die Abbildung 1 zeigt den Aufbau des 100-Gigabit-Testbeds. Für die näherungsweise 60 km lange Dark-Fibre-Verbindung zwischen Dresden und Freiberg kamen die bereits vorhandenen Leitungen der T-Systems zum Einsatz.

Der Alcatel-Lucent 1830 Photonic Service Switch [Alc10a] ist in der Lage 100 Gbit/s auf einer Wellenlänge mittels eines einzelnen Trägers zu übertragen. Die Schnittstelle zwischen der optischen Verbindung und den Netzwerkadaptern bildet ein 7750 SR-12 Ser-

vice Router [Alc10b] von Alcatel-Lucent. Aktuell enthält der Router an beiden Standorten einen Media Dependent Adapter (MDA) mit 1x100 Gbit/s, zwei Adapter mit 5x10 Gbit/s sowie fünf MDAs mit 2x10 Gbit/s. Zusätzlich stehen noch 20x1 Gbit/s Ports zur Verfügung.

Über die verfügbaren 10 Gbit/s Ports sind 18 ProLiant DL160 G6 Server [Hew10] der Firma Hewlett-Packard angeschlossen. Enthalten ist jeweils ein Intel Xeon mit einer Taktfrequenz von 2.67 GHz, 24 GByte Hauptspeicher, ein HP NC550SFP 10GbE Server Adapter, ein HP 4x QDR ConnectX-2 InfiniBand Host Channel Adapter, sowie ein HP StorageWorks 81Q 8Gb Fibre Channel Host Bus Adapter. Zwei der HP Server sind für andere Teilprojekte reserviert, so dass wie in Abbildung 1 dargestellt, in Dresden und Freiberg 16 Knoten für die Durchführung der synthetischen Lasttests zur Verfügung stehen.

Die Speicherlösung in Freiberg ist eine SFA10K [Dat10b] der Firma Direct Data Networks (DDN) mit einer Leistung von 10 GByte/s (Herstellerangabe). Das System in Dresden besteht aus zwei S2A 9900 [Dat10a] mit jeweils 6 GByte/s. Diese Systeme sind über 8Gb Fibre Channel an die HP-Server angebunden. Insgesamt sind 1,5 TByte Speicherplatz verfügbar.

Die InfiniBand Karten der HP-Server (nicht in der Abbildung 1 enthalten) dienen zur Anbindung der lokalen Cluster. Diese wurden als Clients für die I/O-Tests zur Durchsatzoptimierung und maximalen Auslastung des 100-Gigabit-Links verwendet (siehe [Uni10]), wobei die HP-Server als Lustre-Router agieren. Außerdem kommunizieren die HP-Server als Object Storage Targets (OST) des Lustre Dateisystems über die InfiniBand Karten mit dem Metadaten Server (MDS), welcher auf einem extra System installiert ist. In der aktuellen Arbeit wurden die Experimente direkt auf den ProLiant Servern ohne Nutzung der Cluster ausgeführt.

3 Beschreibung der durchgeführten Tests

Um das Potenzial der neu eingeführten 100-Gigabit-Technologie zu quantifizieren werden synthetische Lasttests auf dem Testbed durchgeführt. Der Verwaltungsaufwand beeinflusst dabei die Ergebnisse. Da dieser auf Anwendungsebene häufig nicht exakt definiert werden kann, werden außerdem Messungen auf der Transportebene durchgeführt. In diesem Kapitel wird nach der Definition der Projektziele auf die verschiedenen Testszenarien eingegangen.

3.1 Definition der Projektziele

Neben den klassischen Latenzzeitmessungen und der Bestimmung des Datendurchsatzes mit Hilfe von kontinuierlichen Datenströmen wurden in Zusammenarbeit mit Alcatel-Lucent und T-Systems weitere Projektziele definiert. Im Vordergrund stand dabei das Nutzungsverhalten innerhalb einer realen Umgebung auf das 100-Gigabit-Testbed abzubilden. Dazu gehört das Auftreten von unregelmäßigem Datenverkehr, sogenannten Daten-Bursts.

3.2 Burst-Tests

Auf TCP/IP-Ebene steht, soweit bekannt, kein adäquater Benchmark für die Durchführung von Burst-Tests zur Verfügung. Aus diesem Grund war eine Neuimplementierung notwendig. Eingesetzt wurde hierfür BenchIT [GJ03], ein frei verfügbares Framework für Performance-Messungen. Somit war lediglich die Integration der Funktionalität und die Verknüpfung mit den vorhandenen Schnittstellen notwendig. BenchIT verwaltet die gesammelten Leistungsdaten und bietet eine Oberfläche um diese graphisch darzustellen.

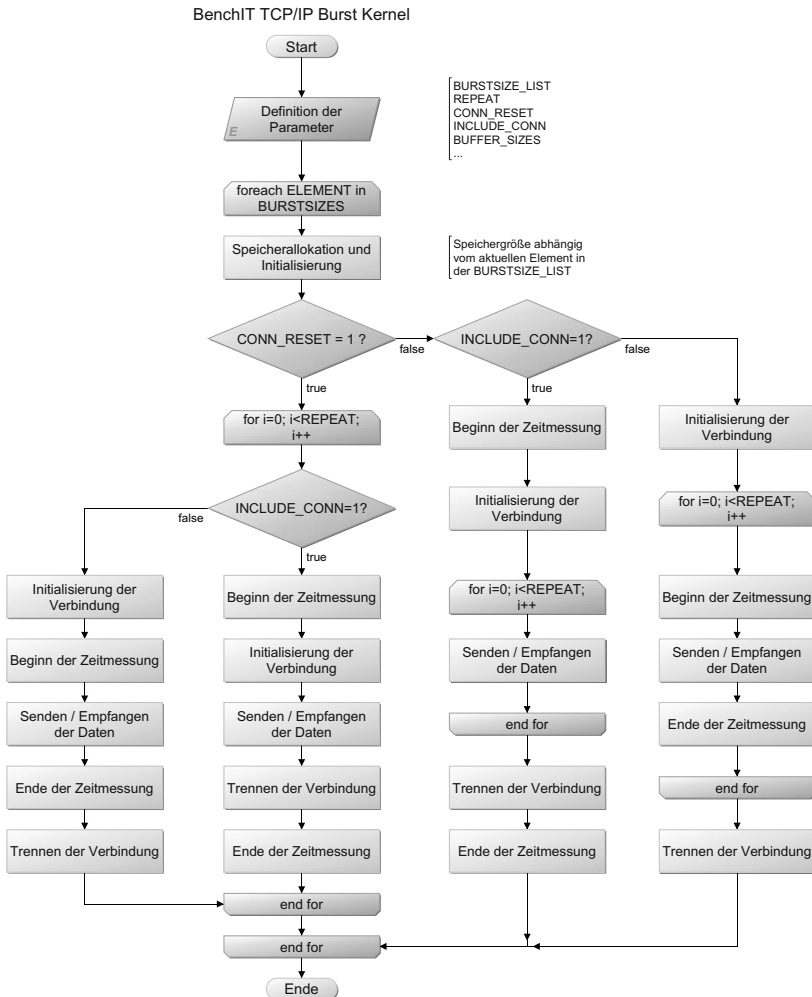


Abbildung 2: Programmablauf des BenchIT TCP/IP Burst Microbenchmarks

Die Abbildung 2 zeigt einen Programmablaufplan des BenchIT-Microbenchmarks. Dabei handelt es sich um ein Client-/Server-Modell wobei der Client zuerst Daten an den Server sendet, welcher diese im Anschluss zurück schickt. Vom Nutzer wird über eine Liste die Größe der Daten-Bursts spezifiziert und wie häufig diese übertragen werden. Zusätzlich können Systemparameter, wie bspw. die Empfangs- und Sendepuffergrößen definiert werden. Diese dienen zur Initialisierung der Verbindung. Weiterhin kann der Nutzer festlegen, ob die Initialisierung der Verbindung zeitlich erfasst wird oder lediglich die Datenübertragung gemessen werden soll. Auch eine Trennung der Verbindung nach jeder Übertragung ist möglich. Unabhängig von den definierten Parametern wird nach jeder Iteration die Differenz aus den beiden Zeitstempeln berechnet und gesichert. Wird die vorgegebene Anzahl an Wiederholungen erreicht, erfolgt die Berechnung des minimalen, maximalen und durchschnittlichen Zeitverbrauchs. Diese Werte werden an BenchIT übergeben und am Ende zusammen mit den Parametern und Umgebungsvariablen in eine Ergebnisdatei geschrieben. Diese wiederum kann im Anschluss mit Hilfe der graphischen Oberfläche von BenchIT visualisiert werden.

Ausgehend von den Burst-Tests auf TCP/IP-Ebene wurden zusätzlich synthetische I/O-Tests auf der Anwendungsebene mit den selben Größen für die Datenbursts und gleicher Anzahl der Prozesse durchgeführt. Dabei kam der Benchmark IOR 2.10.3¹ zum Einsatz. Mit diesem Programm der Universität Kalifornien ist es möglich die Leistung paralleler Dateisysteme zu ermitteln [NLC08]. Zur Synchronisation der Prozesse dient MPI wobei im vorliegenden Testszenario jeder Prozess in eine eigene Datei schreibt. In den Tests wurde ausschließlich POSIX I/O mit dem Flag `O_DIRECT` zur Umgehung der Caches verwendet. Zwischen den einzelnen Tests wurde eine `MPI_Barrier` zur Synchronisation der Prozesse verwendet, um so die Synchronität der Bursts zu gewährleisten. Wie bei den Burst-Tests der Transportschicht wird jeder Test zehnmal durchgeführt. Da die kleinste übertragbare Datenmenge in IOR 8 Byte beträgt, wurden die Größen der Bursts jeweils zum nächst größeren Vielfachen von acht erhöht. IOR unterscheidet zwei für die hier präsentierten Messungen wichtige Werte, die `blockSize` und die `transferSize`. Erstere beschreibt die Größe der Datei und Letztere die Menge der durch einen I/O-Aufruf übertragenen Daten. Um die Vergleichbarkeit der Messungen zu gewährleisten wurden beide Variablen immer auf den gleichen Wert gesetzt wobei allerdings zu beachten ist, dass IOR nur eine Transfergröße von maximal 2 GByte erlaubt. Alle Tests über 2 GByte wurden daher von IOR automatisch in zwei I/O-Operation aufgeteilt. Wie im TCP/IP-Test wurden Daten nur in eine Richtung übertragen. Dies wurde durch reine Schreiboperationen umgesetzt.

Für die IOR Messungen wurde in Freiberg ein Lustre Dateisystem auf den 16 HP-Servern aufgesetzt. Jeder Server verfügt über drei Laufwerke aus dem DDN-System. Jedes Laufwerk wird über InfiniBand als Object Storage Target (OST) an den auf einem eigenen Knoten laufenden Meta Data Server (MDS) exportiert. Somit stehen insgesamt 48 OSTs mit einer theoretischen Bandbreite von 16*8 Gbit/s zur Verfügung. Auf jedem der 16 Server in Dresden wird das Lustre Dateisystem eingebunden. Auf einem weiteren Rechner liegen die Konfigurationsdateien und der IOR Benchmark in einem, über NFS auf allen Servern verfügbaren, Ordner.

¹URL: <http://ior-sio.sourceforge.net>

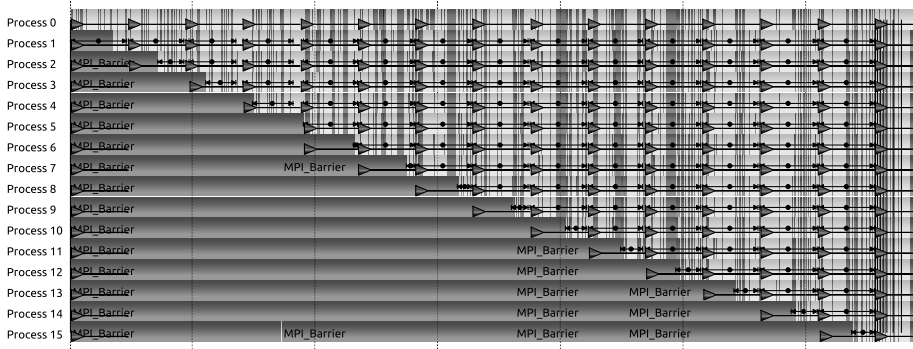


Abbildung 3: Programmspur des Messlaufs visualisiert mit Vampir

Wie im TCP/IP Test werden die Messungen nacheinander mit steigender Anzahl von Rechnern durchgeführt, wobei IOR MPI_Barrier zur Blockierung aller inaktiven Knoten einsetzt. Mit Hilfe des Analysewerkzeugs Vampir [WMM10] wurde eine Programmspur des Testlaufs aufgezeichnet, in welche zusätzlich die Statistikwerte des 7750 Service Routers eingefügt wurden. Der gesamte Messlauf, wie er in Vampir dargestellt wird, ist in Abbildung 3 zu sehen. Dabei ist mit der Farbe rot jegliche Kommunikation mit MPI markiert, gelb stellt I/O-Operationen dar und grün sind Funktionsaufrufe des IOR Benchmarks. Diese dienen hauptsächlich zum Erstellen, Füllen und Leeren der Puffer, welche dann durch die I/O Funktionen auf das Lustre Dateisystem geschrieben werden.

4 Ergebnisse

Im Folgenden werden die in Kapitel 3 beschriebenen Messläufe auf Transport- und Anwendungsebene durchgeführt und diskutiert. Dabei wurde in beiden Fällen über die Anzahl der Knoten und ähnliche Burst-Größen iteriert um die Vergleichbarkeit zu gewährleisten.

4.1 Latenzzeitmessung und Stream-Tests auf der Transportebene

Im ersten Schritt wurde eine Punkt-zu-Punkt-Verbindung zwischen einem Knoten in Dresden und einem Knoten in Freiberg untersucht. Dabei belief sich die Round-Trip-Time (RTT), bestimmt durch die Messung der Zeit für die Hin- und Rückübertragung einer 64 Byte Nachricht, auf 0.72 ms. Diese setzt sich wie folgt zusammen:

- TCP/IP-Stack bzw. Netzwerkkarte: 0.1 ms (Direktverbindung zweier Knoten)
- Alcatel-Lucent 7750 Service Router: 0.08 ms (Messungen an einem Standort)

- Alcatel-Lucent 1830 Photonic Service Switch sowie eine 60 km Glasfaserverbindung: 0.54 ms

Darüber hinaus konnte mit Hilfe einer Path Maximum Transmission Unit (PMTU) von 9000 Byte für die o.g. Verbindung ein Datendurchsatz von 9.92 Gbit/s gemessen werden. Dabei wurde die PMTU durch die in Kapitel 2 beschriebenen Netzwerkkarten definiert. Die Messung erfolgte dabei mit Hilfe des in Kapitel 3 beschriebenen BenchIT-Kernels, wobei anstelle einer Liste von Bursts lediglich ein Wert von 2 GByte übergeben wurde und weder die Initialisierung noch die Trennung der Verbindung gemessen wurden. Die Verwendung höherer Datenmengen zeigte keine besseren Resultate.

Die Ursachen sowohl für die niedrige Latenzzeit, als auch für den hohen Datendurchsatz liegen in der kurzen Übertragungsstrecke von ca. 60 km und in der Tatsache begründet, dass sich alle Knoten in einem Subnetz befinden, wodurch der Routing-Aufwand reduziert wird.

Im nächsten Schritt wurde die Anzahl der Streams, die sich gleichzeitig auf der 100-Gigabit-Strecke befinden, erhöht, indem der für die Punkt-zu-Punkt-Verbindung verwendete Messlauf auf mehrere Knoten verteilt wurde. Dabei konnten mit Hilfe der verfügbaren Monitoring-Funktion des Service Routers die Auslastung des 100-Gigabit-Ports und über BenchIT der Datendurchsatz jeder einzelnen Verbindung aufgezeichnet werden.

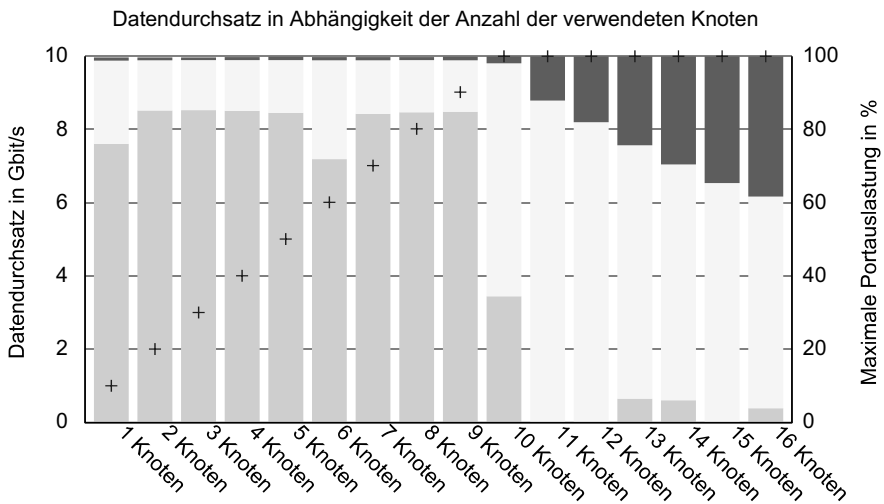


Abbildung 4: Minimaler (grün), durchschnittlicher (gelb) und maximaler (rot) Datendurchsatz, sowie die maximale Auslastung des 100-Gigabit-Ports (blau)

Das Ergebnis in Abbildung 4 zeigt wie die Auslastung des 100-Gigabit-Ports nahezu linear ansteigt, bis sich zehn parallele Streams auf der Leitung befinden. Ab diesem Punkt befindet sich die Strecke unter Überlast, was sich in der steigenden Varianz und dem sinkenden mittleren Datendurchsatz widerspiegelt. Die bereits zuvor sichtbare Differenz zwischen

minimalem und maximalen Datendurchsatz lässt sich mit Hilfe der Initialisierungsphase erklären, in der noch keine volle Datenrate erreicht werden kann.

4.2 Burst-Tests auf der Transportebene

Neben den klassischen Netzwerkparametern wurden, wie in Kapitel 3 beschrieben, auch Burst-Tests durchgeführt. Dabei erfolgte eine nichtlineare, aber reproduzierbare, Iteration über eine Liste, welche Nachrichtengrößen von 1 Byte bis zu 2 GByte enthielt. Die Zeitmessung erfolgte ohne Verbindungsauf- und -abbau und es wurden insgesamt zehn Wiederholung für jede Nachrichtengröße durchgeführt. Aus den daraus generierten Zeitstempeln wurde die minimale, mittlere und maximale Kommunikationszeit bestimmt. Diese Werte wurden zusammengetragen und am Ende des gesamten Messlaufs sortiert ausgegeben. Einen Ausschnitt von 10 MByte bis 2 GByte über 13 Knoten zeigt die Abbildung 5.

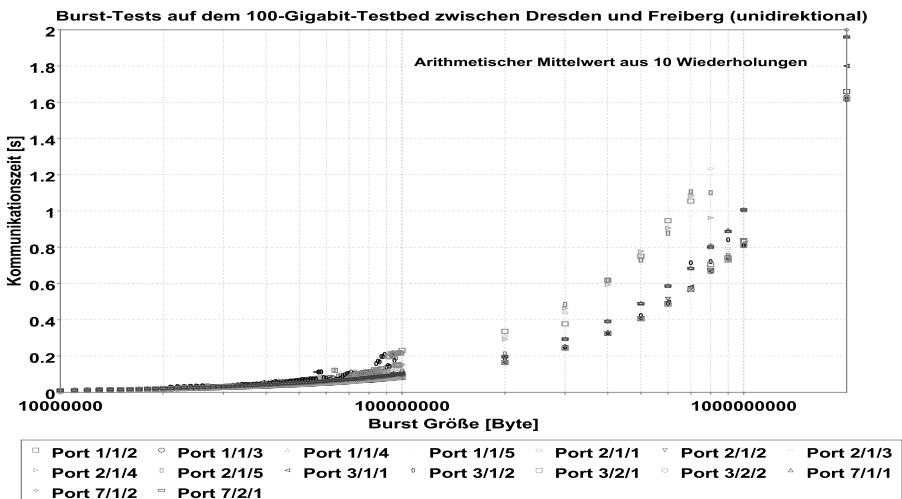


Abbildung 5: Durchschnittliche Kommunikationszeit in Abhängigkeit der übertragenen Datenmenge

Daraus wird ersichtlich, dass bis zu Burst-Größen von 30 MByte kaum Unterschiede in der Kommunikationszeit festzustellen sind, obwohl in Abschnitt 4.1 gezeigt wurde das unter Überlast einzelne Verbindungen benachteiligt werden. Die Erklärung hierfür liegt in der kurzen Übertragungsdauer, so dass die Kapazitäten der Leitung nahezu umgehend wieder zur Verfügung stehen. Bei größeren Datenmengen erhöht sich, wie bereits in Abbildung 4 gezeigt, ebenfalls die Varianz zwischen den einzelnen Verbindungen.

4.3 Burst-Tests auf der Anwendungsebene

Mit IOR wurden insgesamt 2672 Messungen durchgeführt. Die Burst-Größe variierte dabei von 8 Byte bis 4 GByte in unterschiedlichen Schrittgrößen. Es wurde nacheinander für 1-16 Prozesse die verschiedenen Blockgrößen getestet. Dabei wurden 26.173 GiByte an Nutz- und Metadaten von Dresden nach Freiberg gesendet und 75 GiByte an Metadaten in Dresden empfangen. Mit einem Prozess wurden im Durchschnitt 386 MiByte/s übertragen (maximal 468 MiByte/s). Mit vier Prozessen sank die durchschnittliche Rate über alle Burst-Größen auf 350 MiByte/s. Mit allen 16 Prozessen wurden durchschnittliche Datenraten von 280 MiByte/s (Prozess 1) bis 318 MiByte/s (Prozess 12) erzielt. Dass in den Messungen mit allen 16 Servern gerade jene Prozesse, welche in allen Messungen beteiligt waren, die schlechtesten Werte erzielten, lässt sich mit dem *fairness control* Modell von Lustre erklären ([FW09]) welches den Knoten mit weniger Transfervolumen eine höhere Priorität einräumt und somit für eine gleichbleibende Gesamtbandbreite sorgt. Über den kompletten Messzeitraum ergeben sich die in den Abbildungen 6 und 7 dargestellten Werte. In einer Kollaboration mit der Indiana University konnte gezeigt werden,

	240 MiB/s	160 MiB/s	80 MiB/s	0 MiB/s
339.457 MiB/s	/lustre_freiberg/...ior/data.00000000			
334.564 MiB/s	/lustre_freiberg/...ior/data.00000001			
325.524 MiB/s	/lustre_freiberg/...ior/data.00000002			
319.548 MiB/s	/lustre_freiberg/...ior/data.00000003			
316.651 MiB/s	/lustre_freiberg/...ior/data.00000008			
315.753 MiB/s	/lustre_freiberg/...ior/data.00000005			
315.415 MiB/s	/lustre_freiberg/...ior/data.00000004			
315.251 MiB/s	/lustre_freiberg/...ior/data.00000012			
315.199 MiB/s	/lustre_freiberg/...ior/data.00000006			
314.598 MiB/s	/lustre_freiberg/...ior/data.00000011			
314.113 MiB/s	/lustre_freiberg/...ior/data.00000009			
313.697 MiB/s	/lustre_freiberg/...ior/data.00000007			
313.672 MiB/s	/lustre_freiberg/...ior/data.00000010			
310.866 MiB/s	/lustre_freiberg/...ior/data.00000013			
308.489 MiB/s	/lustre_freiberg/...ior/data.00000014			
304.513 MiB/s	/lustre_freiberg/...ior/data.00000015			

Abbildung 6: Durchschnittliche Datenraten für den gesamten Messzeitraum, entspricht 5 GiByte/s am 100-Gigabit-Port

	320 MiB/s	160 MiB/s	0 MiB/s
468.295 MiB/s	/lustre_freiberg/...ior/data.00000000		
466.774 MiB/s	/lustre_freiberg/...ior/data.00000007		
466.189 MiB/s	/lustre_freiberg/...ior/data.00000001		
465.698 MiB/s	/lustre_freiberg/...ior/data.00000008		
465.43 MiB/s	/lustre_freiberg/...ior/data.00000013		
465.042 MiB/s	/lustre_freiberg/...ior/data.00000006		
464.269 MiB/s	/lustre_freiberg/...ior/data.00000005		
462.722 MiB/s	/lustre_freiberg/...ior/data.00000011		
462.704 MiB/s	/lustre_freiberg/...ior/data.00000004		
462.591 MiB/s	/lustre_freiberg/...ior/data.00000009		
462.507 MiB/s	/lustre_freiberg/...ior/data.00000002		
462.435 MiB/s	/lustre_freiberg/...ior/data.00000003		
462.406 MiB/s	/lustre_freiberg/...ior/data.00000012		
461.167 MiB/s	/lustre_freiberg/...ior/data.00000014		
460.017 MiB/s	/lustre_freiberg/...ior/data.00000010		
459.827 MiB/s	/lustre_freiberg/...ior/data.00000015		

Abbildung 7: Maximale Datenraten für den gesamten Messzeitraum, entspricht 7,2 GiByte/s am 100-Gigabit-Port

dass auf der Lustre Netzwerk Schicht, d.h. ohne I/O-Subsystem, *lnet* 94.4% der 100Gbit erreichbar sind [Uni10], dies entspricht 11,79 GByte/s. In einer hoch optimierten Messung konnte mit dem IOR Benchmark beim Schreiben von Dresden auf das Freiburger SFA10K im Spitzenwert 10,8 GByte/s übertragen werden. Dabei wurden allerdings die Lustre Parameter optimiert, ein passendes LFS Striping über alle 3*16 OST's verwendet und der Benchmark selbst lief auf einem über 40Gbit QDR InfiniBand Links an die lokalen 16 HP-Server angeschlossenen Cluster mit bis zu 720 IOR Clients. Die 16 HP Server wurden in dieser Konfiguration als reine *lnet* Router eingesetzt.

5 Zusammenfassung & Ausblick

Die hier vorgestellten Ergebnisse zeigen bereits das Potential der neu eingeführten 100-Gigabit-Technologie. Daraus ergeben sich eine Vielzahl von Anwendungsmöglichkeiten, zum Beispiel die Durchführung dezentraler Backups über große Entfernungen oder die Verbindung von Rechenzentren zur gemeinsamen Nutzung. Dieses Prinzip ist bereits beim Norddeutschen Verbund für Hoch- und Höchstleistungsrechnen (HLRN) mit einer geringeren Übertragungsrate in Anwendung.

Anfang 2011 fand eine Streckenverlängerung auf 200 und 400 km statt. Die dadurch steigenden Latenzzeiten machten eine weitere TCP-Parameteroptimierung erforderlich, um nach wie vor eine volle Linkauslastung zu erreichen. Auch bei den I/O-Messungen waren die bei 60 km im Vergleich zur IO-Hardware vernachlässigbaren Netzwerklatenzen nun ein signifikanter Einfluss. Diese Arbeit diene zur Vorbereitung der Tests auf den längeren Strecken und soll einen Vergleich zwischen den verschiedenen Konfigurationen ermöglichen.

Literatur

- [Alc10a] Alcatel-Lucent. *Alcatel-Lucent 1830 PSS (PSS-32 and PSS-16) - Photonic Service Switch*, 2010.
- [Alc10b] Alcatel-Lucent. *Alcatel-Lucent 7750 SR - Service Router*, 2010.
- [Dat10a] Data Direct Networks, <http://ddn.com/9900>. *Data Direct Networks Silicon Storage Architecture S2A9900*, 2010.
- [Dat10b] Data Direct Networks, <http://ddn.com/10000>. *Data Direct Networks Storage Fusion Architecture SFA10000*, 2010.
- [FW09] Galen Shipman Feiyi Wang, Sarp Oral. Understanding Lustre Filesystem Internals. *Oak Ridge National Laboratory Reporter*, TM-2009(117), 2009.
- [GJ03] Stefan Börner et al. Guido Juckeland. BenchIT - Performance Measurement and Comparison for Scientific Applications. In *PARCO2003 Proceedings*, 2003.
- [Hew10] Hewlett-Packard. *HP ProLiant DL160 G6 Server - Datasheet*, Feb 2010.
- [NLC08] Arifa Nisar, Wei-keng Liao und Alok Choudhary. Scaling parallel I/O performance through I/O delegate and caching system. In *Proceedings of the 2008 ACM/IEEE conference on Supercomputing*, SC '08, pages 9:1–9:12, Piscataway, NJ, USA, 2008. IEEE Press.
- [Uni10] Indiana University. Indiana University Announces Saturation of First Commercial 100 Gigabit Link. <http://www.hpcwire.com/offthewire/Indiana-University-Announces-Saturation-of-First-Commercial-100-Gigabit-Link-108251309.html>, 11 2010.
- [WMM10] Thomas William, Hartmut Mix und Bernd Mohr. Enhanced Performance Analysis of Multi-Core Applications with an Integrated Tool-Chain. In B. Chapman et al., Herausgeber, *Parallel Computing: From Multicores and GPU's to Petascale*, volume 19, pages 701–708. IOS Press, 2010.

ITC Management

Netzwerkboot über Weitverkehrsnetze

Ansatz zur zentralen Steuerung und Verwaltung von Betriebssystemen über das Internet

Sebastian Schmelzer*, Dirk von Suchodoletz*, Gerhard Schneider*

{sebastian.schmelzer,dirk.von.suchodoletz,gerhard.schneider}@rz.uni-freiburg.de

*Albert-Ludwigs-Universität Freiburg

Abstract: Die schnelle Inbetriebnahme von Rechnern, die flexible Nutzung von Maschinen mit verschiedenen Betriebssystemen und Softwareausstattungen sind selbst nach vielen Jahren noch nicht optimal gelöst. Hierfür bietet ein netzwerkbasierter Service mit einer breiten Auswahl zu startender Systeme und Installationsassistenten Abhilfe. Schnelle Datennetze, ein kompaktes Quickstart-OS auf Linux-Basis, ein zentraler Konfigurationsservice und geeignete Protokolle erlauben den einfachen Betrieb von Standardsystemen sowie die Erledigung von Wartungsaufgaben, wie Maschinen-Checks zur Fehlersuche, signifikant erleichtern. Es erfolgt eine Konzentration der Systemlogik auf einen zentralen Dienst, der je nach Anforderung mehrfach redundant ausgelegt werden kann.

Das Konzept basiert dabei auf dem bewährten Netzwerk-Booten, das für einen universelleren Einsatz verbessert wird. Hierzu werden einige Erweiterungen vorgeschlagen, die das Setup und die Konfiguration neuer und in Betrieb befindlicher Maschinen vereinfachen und in bestehende Infrastrukturen einbinden. Dabei wird die Beschränkung der engen Verknüpfung von PXE/DHCP/TFTP aufgehoben. Das erlaubt Rechner aus quasi beliebigen Netzen von einer Reihe verschiedener Bootmedien zu starten. Eine solche Minimalbootumgebung ließe sich für zukünftige Systeme vergleichbar zu den Fastboot-Installationen einiger Hardwarehersteller als Grundausstattung neuer Computergenerationen vorstellen. Auf der Netzwerkkseite könnten der DFN oder einzelne Rechenzentren als Anbieter von neuen Diensten auftreten, die das Serviceangebot erweitern.

1 Motivation

Die Administration von Computern ist trotz aktueller Entwicklungen im Bereich des Software-Deployments immer noch aufwändig, erfordert eine Menge repetitiver Handlungen und bindet teures Personal. Vielfach dauert es sehr lange bis neu beschaffte Maschinen tatsächlich zum Einsatz kommen, was je nach Anzahl eine nicht unerhebliche Verschwendung von Ressourcen darstellt. Dabei unterscheiden sich die Problemstellungen zwischen verschiedenen Institutionen und Körperschaften nur unerheblich, so dass ein generischer Ansatz sinnvoll erscheint. Dieses Paper diskutiert für eine Reihe von Aufgaben ein Konzept, das sich die aktuellen und bewährten Trends der IT- und Netzwerk-Entwicklung zunutze macht. Aktuelle Trends zeigen in Richtung Zentralisierung: Home- und gemeinsame Verzeichnisse sowie die Authentifizierung werden dank schneller Netze überwiegend zen-

tral angeboten. Die konsequente Fortführung erlaubt eine Reihe von Automatisierungen:

- Schnelles Einbinden von neuen Maschinen in die bestehende Infrastruktur, Auswahl einer größeren Palette verschiedener (Linux-)Betriebssysteme für den Einsatz von Desktops, bis hin zu spezialisierten Cluster-Knoten im Grid-Betrieb
- Bereitstellung eines Kiosksystems oder Universaldesktops für Konferenzen, Tagungen oder Gastwissenschaftler
- Anbieten eines "Notsystems", wenn das eigentliche Betriebssystem von Nutzern nicht mehr korrekt arbeitet, was ein gewisses Weiterarbeiten erlaubt und eine Rettungsumgebung bereitstellt
- Erleichterung der Administration von großen Rechnernetzen durch das Anbieten von Spezialsystemen zur Datenrettung, Virussuche, Inventarisierung, der netzwerkgestützten Installation von Linux- und Windowssystemen oder der Bereitstellung einfacher Test- und Evaluationsumgebungen zur Evaluation von Hardware, dem Ausprobieren neuer Betriebssystem- und Softwarevarianten

Ausgehend von einem erprobten Remote-Boot sollen diese Dienstleistungen durch einen OS-on-Demand-Service bereitgestellt werden. Das vorgestellte System lässt sich sowohl lokal innerhalb einer Institution als auch als zentraler Dienst, der von mehreren Universitäten gemeinsam genutzt wird, realisieren. Das Modell ist auf verteilte Rollen angelegt, die es einzelnen Administratoren und Nutzern erlauben ausgehend vom Basis-System weitergehende Anpassungen vorzunehmen ohne dafür jedoch jede Maschine wieder komplett administrieren zu müssen.

2 Booten aus dem Netzwerk

Das Betriebssystem von der konkreten Maschine zu lösen, statt es auf der Festplatte zu installieren, erlaubt einen schnellen Austausch der Hardware nach Ausfällen, Bereitstellung einfacher Test- und Evaluationsumgebungen zur Bewertung der Maschinen-Hardware, dem Ausprobieren neuer Betriebssystem- und Softwarevarianten, den einfachen Umzug von Arbeitsplätzen oder die flexible Nutzung eines Computers mit verschiedenen Betriebssystemen. Rechner via Datennetz zu booten funktioniert seit geraumer Zeit und wird an der Universität Freiburg seit einigen Jahren für eine Vielzahl von Rechnern eingesetzt. Das ursprüngliche Remote-Boot wurde hierzu in den letzten Jahren am Rechenzentrum der Universität Freiburg weiterentwickelt und mit den Möglichkeiten der Virtualisierung speziell für Computer-Pools mit ihren vielfältigen Anforderungen optimiert¹. Es entstand das Open Source Softwarepaket `OpenSLX`, das eine allgemeine Abstraktionsschicht für das Remote-Boot aus dem Netzwerk zur Verfügung stellt². Dieses löst die bisher nur zu Teilen implementierten und sehr unterschiedlichen Ansätze der Linux-Distributionen für Netzwerkstarts oder Live-CDs ab. Gleichzeitig wird die Konfiguration der einzelnen Maschinen von der allgemeinen Bereitstellung eines gemeinsamen Root-Dateisystems für ein

¹Siehe hierzu <http://www.ks.uni-freiburg.de/projekte/ldc>

²OpenSLX-Projektseite: <http://openslx.org>

bestimmtes Linux getrennt. Auf diese Weise arbeiten alle Client-Rechner "stateless" und legen keine maschinenspezifischen Daten im lokalen Dateisystem ab.

Hierzu stellt OpenSLX eine Skriptsammlung bereit, die viele Standard-Distributionen – deren Auswahl sich mit überschaubarem Aufwand erweitern lässt – für den Stateless-Betrieb vorbereitet. Es bietet weiterhin eine Boot-Middleware, die für den Benutzer der laufenden Maschine optimalerweise unsichtbar bleibt. Das jeweilige typische Verhalten der gewählten Distribution bleibt erhalten, die Anpassungen und Änderungen finden im Hintergrund, hauptsächlich während des Bootvorgangs statt. Dieser wurde dabei so optimiert, dass er bestenfalls weniger Zeit als bei einer festplattenbasierten Installation benötigt.

OpenSLX definiert vier Stadien (Abbildung 1), welche die Schritte zur Einrichtung der Maschine sowohl auf dem Boot-Server als auch später auf den einzelnen Client-Rechner bezeichnen. STAGE1 beschreibt die primäre Installation einer Linux-Distribution durch Bootstraps oder durch das Abbild einer Referenzinstallation auf einem weiteren Rechner oder einer virtuellen Maschine in ein Unterverzeichnis des Servers. Dieser Grundinstallation können lokale Anpassungen des Administrators oder Zusatzkomponenten hinzugefügt werden. Im anschließenden Schritt, dem zweiten Stadium, werden hieraus Root-Dateisystem-Exporte erzeugt. Dabei kann es sich um ein via NFS angebotenes Unterverzeichnis oder einen per Blockdevice bereitgestellten SquashFS-Container handeln. Beides erfolgt serverseitig zur primären Einrichtung und nach einem Update oder Änderungen des Root-Filesystems der Clients. Das Ganze ist so angelegt, dass sich auf einem Server theoretisch beliebig viele solcher Installationen anbieten lassen. Alle vom Administrator eingetragenen Exporte stehen sodann als Auswahlliste in einem PXE-Menü³ bereit. Wobei sich selbstverständlich festlegen lässt, welches System nach einem definierten Timeout automatisch startet.

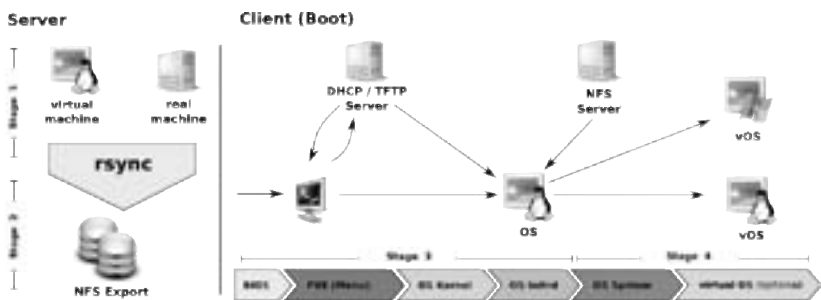


Abbildung 1: Stadien des klassischen OpenSLX

STAGE3 und 4 werden auf dem Client ausgeführt. Hierzu wird das von allen neueren Linux-Kernels unterstützte Konzept der *initramfs* genutzt. In diesem erfolgt das individuelle Setup des einzelnen Rechners durch von OpenSLX erzeugten Skripten. In STAGE3 läuft die individuelle Einstellung der Maschine ab und es werden die vom Administrator bestimmten lokalen Erweiterungen installiert. Im nächsten Stadium übergibt das OpenSLX-Init an das jeweils angepasste Runlevel-System der Distribution.

³PXELinux ist Teil des Syslinux-Pakets von H. P. Anvin: <http://syslinux.zytor.com>

2.1 Vom LAN-Boot ins Weitverkehrsnetz

Die meisten Remote-Boot-Lösungen spielen sich im LAN ab und setzen hierfür auf die weitverbreitete Pre-Boot-Extension (PXE). Diese setzt jedoch die Kontrolle über den lokalen DHCP-Server sowie einen TFTP-Server voraus. Das bedeutet oft eine erhebliche Einschränkung Maschinen in beliebigen Subnetzen aus dem Netz zu installieren oder zu starten. Die Einstellungen auf Subnetzebene mittels DHCP und TFTP-Konfiguration festzulegen, stellt eine nur unbefriedigende Lösung dar. Stattdessen wünscht man sich einen abstrakteren Ansatz, der Konfigurationen an zentralerer Stelle erlaubt. Die Entwicklungen im Bereich des Internets, des DFNs aber auch der Privatanlüsse mit VDSL oder Kabel heben den alten Gegensatz von LAN und WAN zunehmend auf, so dass man zusätzlich zu den dargestellten klassischen LAN-Szenarien Rechner direkt aus dem Internet booten möchte.

Der Ansatz des PXE-Boots lässt sich jedoch nur bedingt auf das WAN transferieren. Oftmals hat der potentielle Benutzer gar kein Zugriff auf die notwendige Konfiguration des lokalen DHCP Servers – entweder aufgrund der Administrationsstruktur oder aber es fehlen die Einstellmöglichkeiten, wie beispielsweise bei den meisten Home-DSL-Routern. Um diese Problematik zu umgehen und von der Abhängigkeit über die Kontrolle von DHCP- und TFTP-Server zu lösen, lässt sich die Linux-Anwendung *kexec* verwenden. Sie erlaubt das Booten eines Betriebssystems oder bestimmter anderer Bootloader, wie *grub4dos* aus dem laufenden Betriebssystem. Das kann nun dazu verwendet werden die Aufgabe von PXE zu ersetzen. Es bedarf zu Beginn eines (minimalen) Linux-Systems mit Netzwerkunterstützung, dann können Kernel und die von OpenSLX generierte *initramfs* über das Internet heruntergeladen werden und mit *kexec* gestartet werden [SvS09].

2.2 WAN-Boot: Prototyp auf Basis des klassischen OpenSLX

Der eben vorgestellte Ansatz schafft eine Abhängigkeit: die eines laufenden Linux-Systems. Daher muss das zugrunde liegende Mini-Linux-System klein und flexibel sein. Es zeigt sich, dass sich ein solches System bestehend aus einem Kernel mit einer möglichst breiten Palette an unterstützten Netzwerkkarten sowie eines *initramfs* mit *kexec* in einem wenige Megabyte großen Image unterbringen lassen. In einem ersten Prototyp wurde das Erstellen der Bootmedien (CD-Image oder USB-Bootsektor) mit der Hilfe einer Web-Applikation, dem Pre-Boot-Server (PBS) erstellt [Sch10]. Dieser diente gleichzeitig als Gegenstelle für die auf dem Bootmedium untergebrachten Skripte. Zudem wurde OpenSLX so modifiziert, dass es bei der Erstellung der PXE-Menüs zusätzlich die Informationen über die bootbaren Systeme an die Web Applikation übergeben hat. Aus diesen Daten ließen sich dann in der Verwaltungsoberfläche des PBS Menüs zusammenstellen und auf einzelne Bootmedien abbilden. Zusätzlich konnte die Zuordnung der Menüs Abhängig von IP-Adressräumen gemacht werden. Für die Auslieferung des Kernels und der von OpenSLX erzeugten *initramfs* diente der PBS als HTTP-Schnittstelle zu dem lokalen TFTP-Server.

Mit dem PBS ließ sich ein erstes lauffähiges Model eines WAN-Boot Dienstes betreiben,

das jedoch Schwächen offenbarte. Das klassische OpenSLX entsprach nicht den Anforderungen eines flexiblen, vielseitig konfigurierbaren Dienstes.

3 Konzept: OS-on-Demand

Aus den Erfahrungen mit der auf OpenSLX basierenden Lehrpool-Umgebung an der Uni Freiburg und einem ersten WAN-Boot-Prototypen wird im Moment die nächste Generation des OpenSLX entwickelt, die erweiterten Anforderungen gerecht werden soll. Deshalb erfolgt eine logische Trennung von Maschinenkonfiguration und Bereitstellung der Root-Filesysteme für die verschiedenen Linux-Varianten. Die Konfiguration wird auf einem Server konzentriert und den lokalen Administratoren mit Filesystem-Overlays weitgehende Anpassungsmöglichkeiten eingeräumt.

Der bisherige Ansatz von OpenSLX sah vor, dass es in der Regel eine kleine Gruppe von Administratoren gibt, die selten die Konfiguration der Systeme ändern. Um jedoch einen über das lokale Netz hinausgehenden Dienst anbieten zu können, müssen die Komponenten zur Konfiguration der Systeme neu überdacht werden.

Bislang lag der Fokus der Entwicklung von OpenSLX in der Bereitstellung einer wartungsarmen, flexiblen Computerpool-Umgebung. Dies bleibt im neuen Konzept erhalten. Zudem sollte es möglich sein den Dienst einfach in bestehende Infrastrukturen einzubinden. So sollen beispielsweise die Anbindung an die institutseigene Userauthentifizierung und die Einbindung der Heimatverzeichnisse der Benutzer funktionieren, ebenso wie die Erweiterung der Betriebssystemvorlagen durch zusätzliche Filesystemschichten (Overlays) mit vom Standard abweichenden Softwarepaketen. Dieses macht eine Mandantenfähigkeit der Verwaltungsoberfläche notwendig, damit die Administratoren der Institution in der Lage sind, Einstellungen und Erweiterungen für ihren Zuständigkeitsbereich eigenständig vornehmen zu können.

Konfigurierbarkeit für den Einzelanwender Der Einzelanwender soll die Möglichkeit haben, schnell und einfach "Live Systeme" auf seinem Rechner zu nutzen⁴. Zudem soll dem Benutzer gestattet werden eigene Einstellungen vorzunehmen, wie beispielsweise das Setzen eines Standardsystems beim Booten, das Einbinden von lokalen Speichermedien oder das Einrichten eines Autologins einer definierten X-Session. Für diesen Zweck soll es für Nutzer erlaubt sein, sich am Dienst zu registrieren, um die Einstellungen an zentraler Stelle zu hinterlegen.

Modifikationen der ursprünglichen Linux Distributionen Die notwendigen Veränderungen, die für die Vorbereitung einer üblichen Linux Installation zum Booten über das Netzwerk notwendig sind, haben sich seit dem Beginn der Entwicklung von OpenSLX massiv verändert. So muss ein Großteil der Hardwarekonfiguration nicht mehr manuell vorgenommen werden, da das System diese automatisch erkennt und einrichtet – so be-

⁴Vergleichbar mit dem Dienst *boot.kernel.org*

darf es beispielsweise im Normalfall keiner Konfiguration der X Windows Oberfläche mehr. Auf der anderen Seite gibt es neue Anforderungen, um ein System über das Netzwerk booten zu können. Die meisten der aktuellen Linux Distributionen optimieren den Startprozess ihrer Systeme massiv, beispielsweise durch die Verwendung von parallelisierten Init Systemen (z.B. *upstart*). Dies erschwert jedoch grundlegend die Integration des Netzwerk-Boots. Der Eingriff in das Init-System des zu bootenden Betriebssystems ist notwendig, da Aufgaben, wie das Netzwerksetup und das Einbinden des Root-Filesystems, die normalerweise das Betriebssystem eigene Init selbst ausführen würde, bereits durch die für den Netzwerk-Boot optimierte *initramfs* ausgeführt wurden. Die meisten Modifikationen, die nötig sind, um ein System für den Netzwerk-Boot vorzubereiten, sind also nicht mehr Hardwareabhängig sondern Distributionsabhängig. Dies erlaubt die Modifikationen nicht wie bislang im von OpenSLX generierten *initramfs* auszuführen, sondern auf dem Server bei Bedarf zu erzeugen und für wiederholte Nutzung zwischenspeichern.

3.1 Komponenten

Aus den soeben beschriebenen Anforderungen ergibt sich der im Folgenden beschriebene Entwurf des OS-on-Demand(OSoD) Dienstes. Bevor auf diesen genauer eingegangen wird, werden die beteiligten Komponenten definiert:

PBL steht für Pre-Boot-Linux, es handelt sich hierbei um eine bereits in einer ähnlichen Form erprobte basierende Mini-Linux Umgebung. Sie hat die Aufgabe, die Funktion von PXE zu ersetzen und erweitern. Sie wird auf einem beliebigen Bootmedium ausgeliefert - das heißt sie kann auf USB-Sticks, CDs, Festplatten oder wiederum über PXE bereitgestellt werden. Das Bootmedium umfasst einen kompakten Linux Kernel, mit Unterstützung für die meisten Netzwerkkarten und der einer auf Busybox basierenden *initialramfs*. Neben den Standardkommandos zum Einrichten des Netzwerks und Nachladen weiterer Komponenten wird das PBL auch eine auf dem Linux-Framebuffer laufenden Client Applikation für die Interaktion mit dem Benutzer enthalten, welche direkt mit dem OSConfig-Server kommuniziert.

Der **OSConfig-Server** ist eine Web Applikation. Sie lässt sich in mehrere Komponenten unterteilen. Zum einen gibt es eine Komponente, die der Verwaltung des kompletten Systems dient. Sie beinhaltet das Rechtekmanagement und verwaltet die zugehörigen OSRootfs Server/Proxies, die bootbaren Betriebssystem-Vorlagen als auch lokale Overlays – Erweiterungen zu den Basis-Betriebssystem-Vorlagen; zudem lassen sich hier neue Bootmedien erstellen. Eine weitere Komponente verwaltet die Einstellungen von Einzelbenutzern. Auf diese Komponente kann sowohl über ein Web Interface zugegriffen werden als auch über die im PBL enthaltene Client Applikation. Die letzte und entscheidende Komponente ist ein Interface zur Erzeugung und Verteilung von Kernel, Initrd und der Konfigurationslayer, die aus einer gewöhnlichen Linux-Installation ein über das Netz bootbares System machen. Diese Konfigurationen werden im *initialramfs* des zu bootenden Systems abgerufen und über das nur lesbar eingebundene Root-Filesystem auf einen beschreibbaren

Layer (Unionfs⁵/Aufs⁶) gelegt.

Die **OSRootfs-Server** bzw. **-Proxies** dienen dazu, aus bestehenden Linux-Installationen Vorlagen-Images für den Netzwerk-Boot zu erzeugen. Diese werden mittels `rsync` kopiert, wobei gleichzeitig eine Vielzahl von laufzeitspezifischen Dateien bereits aussortiert werden, wie Log- oder Cache-Dateien. Diese "aufbereiteten Vorlagen" werden dann über das Netzwerk via NFS oder ein Network Block Device zur Verfügung gestellt. Die OSRootfs-Server erhalten ein Interface, um Vorlagen Images auszutauschen. Diese Schnittstelle werden auch die sogenannten OSRootfs Proxies, lokale Proxy Server für Vorlagensysteme, nutzen.

3.2 Aufbau

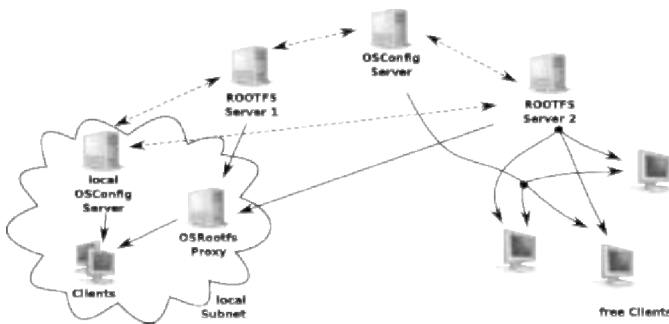


Abbildung 2: Aufbau des OS-on-Demand Dienstes

Das Zusammenwirken der einzelnen Komponenten veranschaulicht Abbildung 2. Prinzipiell lassen sich zwei Anwendergruppen unterscheiden: Institutionen (1) und der Einzelbenutzer (2). Institutionen haben in der Regel eine größere Anzahl von Rechnern zu verwalten und dabei komplexere Anforderungen an die Modifikationen der zu bootenden Linux Vorlagen, wie die Einbindung von Benutzerverzeichnissen und der zentralen Benutzerauthentifizierung. Zu diesem Zweck bietet sich der Betrieb eines eigenen OSConfig Servers an (3). Ein OSConfig Server ist mit einem oder mehreren OSRootfs Servern (5) verknüpft, da der OSConfig Server für die Erzeugung der Netzwerk-Boot Konfiguration Zugriff auf das per NFS vom OSRootfs Server eingebaute Root-Filesystem der zu bootenden Linux-Variante benötigt. Beim Betrieb einer großen Anzahl von Maschinen bietet sich auch die Verwendung eines OSRootfs Proxies (4) an, der alleine der Zwischenspeicherung der genutzten Vorlagensysteme im lokalen Netz dient. Die Einzelbenutzer des Dienstes können sich an einem öffentlichen OSConfig Server (6) registrieren und die von diesem vorgehaltenen Systeme aus dem WAN booten.

⁵Union File System <http://www.unionfs.org>

⁶Another Union File System <http://aufs.sf.net>

3.3 Bootprozess

Der Bootvorgang eines Clients unterscheidet sich an einigen Punkten vom klassischen Hochfahren eines lokalen Systems. Die Abbildung 3 zeigt schematisch sowohl den zeitlichen Ablauf als auch die dabei beteiligten Komponenten. Nach dem Einschalten des Rechners sucht zunächst das BIOS nach einem bootbaren Medium. Hierbei kann für das Booten einer OSoD Session eines der vom OSConfig erzeugten Bootmedien (1) verwendet werden. Es können – sofern vom Rechner unterstützt – optische Datenträger (CD/DVD), Flash-Speicher (SD/CF Karten), USB-Sticks, die lokale Festplatte oder auch PXE dienen. Von diesem Medium wird das Minimal-Linux PBL gestartet. Nach dem Netzwerksetup wählt der Benutzer mittels einer auf dem Linux-Framebuffer basierenden graphischen Benutzeroberfläche das zu ladende Betriebssystem aus, deren Inhalte die Applikation direkt vom OSConfig Server bezieht (2). Zusätzlich können an dieser Stelle Einstellungen geladen und verändert werden. Anschließend werden Kernel und *initramfs* des zu bootenden Systems geholt und mittels *kexec* gestartet. Das Programm *kexec* ist eine kleine Anwendung, die es erlaubt, ein Linuxsystem aus einem laufenden System heraus zu booten, ohne dabei einen Neustart durchführen zu müssen. Es ist jedoch bis auf die Ausnahme der Kernel Command Line (KCL) nicht möglich (Konfigurations-)Daten in das neue System zu übermitteln. Die KCL ist eine längenbeschränkte Zeichenkette, mit der dem Kernel Parameter übergeben werden können. In dieser platziert das PBL einen Sitzungsschlüssel, um zu einem späteren Zeitpunkt die passenden Konfigurationen vom OSConfig Server zu laden (3). Für die Erzeugung angepasster Konfigurationen besitzt der OSConfig Server Zugriff auf den OSRootfs Server (4). Nachdem der OS-Kernel geladen ist, wird die für das zu bootende System erzeugte WAN-Boot-Initramfs geladen. Diese ruft vom OSConfig Server den Konfigurationslayer ab und legt ihn über das nur lesbar vom OSRootfs Server eingebundene Root-Filesystem (5). Optional werden an dieser Stelle noch weitere Schichten in den Dateibaum eingebunden. Mit dem Verlassen des *initramfs* folgt der reguläre Systemstart. Um auch andere (proprietäre) Systeme über das soeben vorgestellte Verfahren laden zu können, ist es vorstellbar, nach dem Boot des zugrunde liegenden Linux Systems noch ein weiteres virtualisiertes Betriebssystem zu starten (6).

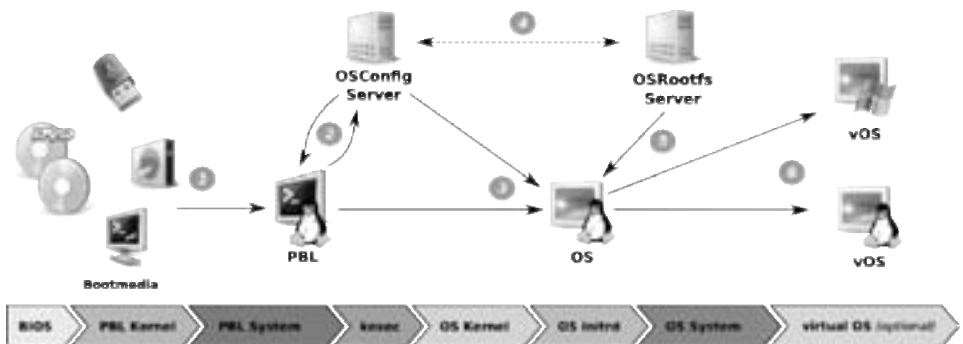


Abbildung 3: Bootprozess eines Clients des OSoD Dienstes

3.4 Verteilung der Root-Filesysteme

Bei der Umsetzung des vorgestellten Konzepts gilt es noch einige Fragen zu beantworten. Die zunächst wichtigste Aufgabe wird die Auswahl des geeigneten Wegs die Distributions-Vorlagen zu verbreiten. Im lokalen Netz hat sich hierfür NFS bewährt, jedoch zeigten die ersten Experimente mit dem Boot aus dem WAN, dass NFS hierbei an seine Grenzen stößt. Vielversprechender waren hierbei die Versuche mit Network Block Devices. Hierbei wird mittels Squashfs ein komprimiertes Loopdevice erzeugt, welches dann über ein entsprechendes Protokoll verteilt wird. Da diese Protokolle nur auf Blockebene arbeiten, besitzen sie wesentlich weniger Overhead. In Frage kommen hierfür das bereits im Kernel enthaltene NBD, sowie die am Lehrstuhl entwickelten Distributed Network Block Devices (DNBD, DNBD2) [Gre07].

4 Anwendungsfälle

Von den in der Motivation vorgestellten Einsatzmöglichkeiten sollen zwei vertieft werden, um die Verwendung des Konzept des OS-on-Demand zu verdeutlichen.

Universaldesktop für Veranstaltungen (Konferenz-Pools) Eine sich regelmäßig wiederholende Aufgabe ist die Bereitstellung von Desktopsystemen für Konferenz-, Tagungs- oder Seminarteilnehmer zum Online-Gehen. Dabei existieren nur wenige Unterschiede, an welcher Universität oder Forschungseinrichtung die Veranstaltung stattfindet. Hier bietet sich der Einsatz des WAN-Boots einer zentral vorgehaltenen Standardumgebung an, die bei Bedarf lokal erweitert wird. Ein wesentlicher Vorteil liegt darin, dass quasi beliebige Maschinen in einen solchen Konferenzpool abgestellt werden können, da keine Änderungen an eventuell existierenden Installationen auf der Festplatte vor oder nach der Veranstaltung vorgenommen werden. Dieses erübrigt auch das Vorhalten spezieller Maschinen, die nur zu diesem Zweck zum Einsatz kommen.

Dynamische Lehrpool-Steuerung Eine wirklich dynamische Steuerung der im Normalfall zu bootenden Betriebssysteme war mit dem klassischen LAN-Boot-Ansatz aufgrund der Vielzahl beteiligter Dienste, die je nach Institution in unterschiedliche Aufgabenbereiche (Administrationsdomänen) fallen, nicht trivial. Im Gegensatz zu den relativ statischen, über TFTP ausgetauschten PXE-Menüs bieten die dynamisch in der Web-Applikation *OS-Config Server* erzeugten Menüs eine Vielzahl von Varianten zur Manipulation. So lässt sich die Auswahl der bootbaren Systeme sowie die Defaulteinträge abhängig von Standort, Hardware, Bootmedium oder auch Uhrzeit steuern. In der denkbaren Kombination mit Techniken wie Wake-On-Lan besitzt der Administrator ein umfangreiches Werkzeug zur Steuerung seiner Pool Umgebung. Dies ermöglicht beispielsweise, dass zu einer bestimmten Uhrzeit ein vordefiniertes System für eine Lehrveranstaltung automatisch startet, dass Dozenten-PCs andere Voreinstellungen erhalten oder aber dass die Pool Rechner in ungenutzten Zeitintervallen andere Aufgaben erfüllen, wie über Nacht ihre Rechenkapazität in ein Cluster einzubinden.

5 Ausblick

Die präsentierte Lösung basiert auf langjährigen Erfahrungen im Bereich des Leih- und Lehr-Pool-, Cluster- und Desktopbetriebs am Rechenzentrum der Universität Freiburg. Während sich in statischen Umgebungen mit sehr fest umrissenen Aufstellungsorten und Maschinen die klassische PXE-Lösung anbietet, stößt sie für größere Installationen mit verteilten Managementrollen für die verschiedenen institutionellen Einheiten an ihre Grenzen. Sollten einzelne Nutzer zusätzlich die Möglichkeit besitzen, ihre Maschinen in einem gewissen Grade selbst zu verwalten, eignet sich ein zentraler Konfigurationsserver.

Der Linux-Diskless-Betrieb und die verteilte Administration des Root-Filesystems sind schon länger erfolgreich getestet. Für das WAN-Boot-Modul und den Konfigurationsserver hingegen existieren bisher nur Prototypen, die im Moment durch ein studentisches Teamprojekt verfeinert werden.

Das vorgeschlagene Betriebsmodell erhebt keinen Ausschließlichkeitsanspruch, sondern ist darauf angelegt mit anderen zu koexistieren. Kein Administrator muss die Hoheit über seinen DHCP-Server aufgeben oder von RIS Abschied nehmen. Ein Nutzer kann jederzeit sein lokal installiertes Betriebssystem booten und den OS-on-Demand Dienst lediglich zur Neuinstallation oder für Hardwaretests verwenden. Auf dem gezeigten Weg könnten die sich teilweise ausschließenden Softwareinstallations- und Managementinfrastrukturen unter einem einzigen Einstiegspunkt vereinigt werden.

Wenn man die Idee der bereits auf dem Mainboard oder Laptop vorinstallierten Minimal-Systeme⁷ konsequent weiterverfolgt, dann könnte man sich in Zukunft ein universelles Mini-Boot-System (PBL) parallel dazu oder als Ersatz vorstellen. Die Platzanforderungen sind mit circa 15 MByte marginal. Ein Update ist nicht notwendig, da die relevante Hardware (Netzwerkkarte) fix ist. Das könnte Administratoren das Leben bei der Erledigung von Standardaufgaben erheblich vereinfachen und das Herumtragen und Aktualisieren von Installations- und Systemrettungsdatenträgern jeder Form überflüssig machen.

Literatur

- [Gre07] Vito Di Leo Gremmelspacher. Development of a Distributed Network Block Device for Unicast Networks. <http://www.ks.uni-freiburg.de/download/bachelorarbeit/SS07/06-07-dnbd2-dileo/ba-dnbd2-dileo.pdf>, 2007.
- [Gre08] Kate Greene. Booten ohne Verzögerung. *Technology Review*, 2008. <http://heise.de/274854>.
- [Sch10] Sebastian Schmelzer. Extending the playground. <http://www.ks.uni-freiburg.de/projekte/ldc/extendingtheplayground.pdf>, 2010.
- [SvS09] Sebastian Schmelzer und Dirk von Suchodoletz. Network Linux Anywhere. *Linux-Tag 2009*, 06 2009. <http://www.ks.uni-freiburg.de/projekte/ldc/networklinuxanywhere-paper.pdf>.

⁷Beispielsweise der Splashtop auf ASUS Geräten, [Gre08]

Eine Taxonomie und Bewertung von Cloud Computing-Diensten aus Sicht der Anwendungsentwickler

Michael Kretzschmar und Mario Golling

Universität der Bundeswehr München

Fakultät für Informatik

D-85577 Neubiberg

Email: {michael.kretzschmar, mario.golling}@unibw.de

Abstract: Die Idee des Cloud Computings ist eine der meistversprochensten der jüngeren IT-Geschichte. Die Kombination aus bekannten Konzepten wie das der verteilten Systeme mit neuen technischen Möglichkeiten auf den Gebieten der Virtualisierung und Datenübertragung ermöglichen eine neue Stufe der Leistungsfähigkeit und Wirtschaftlichkeit. Abseits der unumstrittenen Vorzüge der Cloud gibt es viele offene Fragen auf diesem recht jungen Gebiet der IT. Basierend auf der Tatsache, dass die meisten Anbieter beim Implementieren eigene Normen entwickeln, existieren zur Zeit nur sehr eingeschränkt Standards bzw. Konventionen bezüglich Datenübertragung, Sicherheit und Interoperabilität etc.

Neben Nutzern und Anbietern von Cloud-Diensten gibt es die Anwendungsentwickler. Die Anbieter von Cloud-Diensten stellen Entwicklern Anbindungen unter anderem in Form von Application Programming Interfaces (APIs) zur Verfügung, durch die Programme den jeweiligen Dienst einbinden können. Doch eben diese Anbindungen sind ebenfalls nicht standardisiert und werfen bei Entwicklern grundsätzliche Fragen bezüglich Anbieterabhängigkeit, Interoperabilität, Programmiersprachen- und Protokollauswahl etc. auf. Zudem fehlt es aus Entwicklersicht an Übersichten und Entscheidungshilfen zu den verschiedensten Anwendungsspektren.

Spezifische Antworten auf diese Fragen zu geben ist nur sehr eingeschränkt möglich, da jeder Entwickler eigene Ziele und Voraussetzungen hat, die die Entscheidungen bezüglich einzelner Punkte maßgeblich beeinflussen. Im Sinne einer Entscheidungsunterstützung wird in diesem Beitrag ein Überblick über aktuelle Implementierungen von Cloud-Diensten gegeben. Hierbei wird aufbauend auf einem erweiterten SPI-Modell betrachtet, welche Arten von Diensten gegenwärtig existieren und welche speziellen Verfahren für allgemeine Aktionen, insbesondere den Zugriff auf die Dienste, benutzt werden können. Es werden Optionen gegenübergestellt und deren Implementierungen evaluiert. Dabei spielt der Vergleich von SOAP- mit REST-basierten Anbindungen eine große Rolle. Als Basis für Entwickler wird zusammenfassend zu jedem untersuchten Spektrum (Computing, Speicher, Plattform und Software) eine Entscheidungshilfe mit den jeweiligen Optionen präsentiert.

1 Motivation

In zahlreichen White Papers verschiedener Anbieter von Cloud-Diensten sind die Vorteile des Arbeitens in der Cloud aufgezählt [TSW, GGW09, SFW10]. So sparen beispielsweise Großfirmen auf Grund der geringeren Kosten im Vergleich zur Nutzung/Instandhaltung

eigener Rechenzentren. Insbesondere kleine, mittlere oder neugegründete Unternehmen können sich auf Grund der Philosophie vertragsungebundener Dienstleistungen (*Pay-as-You-Go*) risikofrei IT-Infrastrukturen leisten.

Neben den Endanwendern, die die Cloud verwenden, gibt es die Entwickler, die Software - die wiederum einen Cloud-Dienst als Grundlage nutzt - entwickeln. Hierzu stellen die Anbieter von Cloud-Diensten Entwicklern Anbindungen unter anderem in Form von Application Programming Interfaces (APIs) zur Verfügung, durch die Programme den jeweiligen Dienst einbinden können. Doch eben diese Anbindungen sind ebenfalls nicht standardisiert und werfen bei Entwicklern Fragen auf:

- Welche Programmiersprache soll verwendet werden?
- Welche Standards unterstützt der Cloud-Anbieter?
- Wird das Programm auch mit anderen Anbietern funktionieren?

Im Sinne einer Entscheidungsunterstützung wird daher im Folgenden ein Überblick über aktuelle Implementierungen von Cloud-Diensten gegeben. Dabei wird betrachtet, welche Arten von Diensten gegenwärtig existieren und welche speziellen Verfahren für allgemeine Aktionen, insbesondere den Zugriff auf die Dienste, benutzt werden können.

2 Das erweiterte SPI-Modell

Mit der Entwicklung des Cloud Computings begannen verschiedene Anbieter unterschiedliche Arten von Diensten anzubieten. Bereits sehr schnell deckte das Gesamtportfolio von Cloud-Diensten eine große Spannweite ab. Vom Angebot *einfacher* Dienste, wie der Nutzung von sogenannten *Speicherdiensten*, die im einfachsten Fall mehr oder weniger einer in das Internet ausgelagerten Festplatte gleichen (mit all den Vorteilen, wie weltweiter Verfügbarkeit, redundanter Speicherung etc.) bis hin zu komplexen Customer-Relationship-Management-Software (CRM). Zur besseren Übersicht hat sich sehr schnell ein Service Modell, das sogenannte SPI-Modell (SPI als Abkürzung für die drei Ebenen *Software as a Service*, *Platform as a Service* und *Infrastructure as a Service*) [Bia08] [MKL09], in der Fachschaft durchgesetzt, welches allgemein sämtliche Cloud-Dienste einer von drei aufeinander aufbauenden Ebenen zuteilt:

2.1 Infrastructure as a Service (IaaS)

IaaS umfasst Dienste, die grundlegende Ressourcen zur Verfügung stellen. Dabei kann es sich um Speicherplatz oder Rechenleistung handeln, die meist in Form von virtuellen Servern auf der physischen Hardware der Anbieter laufen und auf die - durch Abstraktion - in Form von Diensten zugegriffen werden kann. Sie kann weiter unterteilt werden in:

Cloud Computing - Der Bereich, der dem großen Konzept seinen Namen gibt, beschreibt

insbesondere Rechenleistung und Arbeitsspeicher. Hierzu werden die überwiegend in Rechenzentren befindlichen Ressourcen virtualisiert. Das heißt im Cloud-Kontext: Zusammenfassung als “großes Ganzes” und danach wieder Aufteilung in einzelne virtuelle Server. Verwendet werden hierzu Virtualisierungssysteme wie Microsofts Hyper-V [War08] oder der Xen Hypervisor von Citrix [Chi07] [Hes10].

Cloud-Speicher - Speicherressourcen lassen sich bei genauer Analyse des Marktes noch einmal in zwei verschiedene Ausprägungen unterteilen:

- *Blockspeicher* - Als Block angebotene persistente Speicherressourcen, die sich - vergleichbar mit Network-Attached Storage (NAS) - im Rechner als Laufwerk einbinden lassen, werden als Blockspeicher bezeichnet. Hierbei ist zu beachten, dass die Größe des Blocks, also die Größe des Laufwerks, im nachhinein nicht mehr verändert werden kann. Ein Beispiel für solch einen Dienst ist Zettas Enterprise Network Storage Service [Zet10].
- *BLOB-Speicher* - Im Vergleich mit Blockspeicher stellt BLOB-Speicher (BLOB als Akronym für Binary Large Object) ein agileres Konzept dar, da die Größe der Speicherressourcen hier praktisch unbeschränkt ist. Aufgrund der Möglichkeit Dateien öffentlich mittels Web-Link publik zu machen, muss in der Regel darauf geachtet werden, dass global einzigartige Namen verwendet werden. Amazons Simple Storage Service (S3) [S310], Googles Storage for Developers [SfD11] sowie Windows Azure Storage [AZS11] sind nur einige Beispiele.

2.2 Platform as a Service (PaaS)

PaaS beschreibt Dienste, die eine Laufzeitumgebung für Anwendungen bereitstellen und erlaubt es dem Verbraucher, eigene Anwendungen oder Webseiten in der Cloud zu betreiben. Auf variable Last der Anwendungen durch Berechnungen oder Zugriffe kann schnell reagiert werden, indem Ressourcen auf IaaS-Ebene flexibel genutzt werden (*automatische Skalierung*). Ein weiterer Vorteil ist, dass hier der Anbieter sicherstellt, dass keinerlei Arbeiten für das Erstellen einer Laufzeitumgebung durchgeführt werden müssen, da z. B. Frameworks, Patches schon installiert sind.

2.3 Software as a Service (SaaS)

SaaS beinhaltet Programme, die auf einer Cloud-Infrastruktur ausgeführt werden und beschreibt in sich geschlossene Anwendungen, die lediglich in einigen vorgegebenen Optionen konfiguriert werden. Dazu zählen Cloud-basierte E-Mail- oder CRM-Dienste wie die Sales Cloud 2 von Salesforce [SFS11]. Hier kann sich der Benutzer mit seinem Browser anmelden und mit dem Sales Cloud-Dienst arbeiten, ohne weitere Software auf seinem Rechner zu installieren oder zu konfigurieren. Sämtliche Abläufe finden in der Cloud statt. Lediglich die Darstellung der Benutzeroberfläche im Browser belastet das System des

Endanwenders. Xignite beispielsweise dient mit Finanzmarktdaten, die Entwickler mittels einer bereitgestellten Anleitung in ihre eigenen Programme integrieren können [Xig10].

3 Zugriffsarten auf Cloud-Dienste

Für den Zugriff auf die jeweiligen Dienste existieren verschiedene Konzepte:

- *Web GUIs* - Intuitive grafische Benutzeroberflächen werden praktisch von sämtlichen Anbietern zur Verfügung gestellt und sind mittels Browser zu bedienen.
- *Remote Administration* - Im Gegensatz zu Web GUIs werden hier Dienste von einem entfernten Ort, ohne wahrnehmbare Zwischenschicht zwischen Benutzer und Dienst, gesteuert. Zur Kommunikation werden beispielsweise Befehle über eine Terminal-Sitzung übertragen.
- *API Anfragen* - Die APIs der Cloud-Dienstanbieter stellen Entwicklern Vorgaben und Beispiele für Anfragen, sogenannte Requests, zur Verfügung, um Ihren Dienst entsprechend nutzen zu können. Die Herausforderung dabei besteht für die Entwickler darin, jeweils abhängig vom konkreten Cloud-Dienst individuell korrekte Anfragen zu generieren.

Der elementare Unterschied im Bereich der API Anfragen ist die Formulierung sowie die Übertragungsart. Cloud-Dienstanbieter stellen APIs bereit, die meist auf REST oder SOAP basierend beschrieben werden. REST ist eine Abkürzung für Representational State Transfer und beschreibt eine Sammlung von Vorgaben, die den Informationsaustausch vereinheitlicht, erleichtert sowie die Leistung dieser Systeme erhöht. So wird durch REST z.B. gefordert, dass jede Anfrage sämtliche Informationen enthält, die vom Server benötigt werden (Stateless) oder dass Informationen, die den Client interessieren könnten, in der Repräsentation der Ressource verfügbar sein sollen. SOAP, anfangs ein Akronym für Simple Object Access Protocol, hingegen ist ein auf XML basierendes Kommunikationsprotokoll für Anwendungen, das seine Nachricht standardisiert verpackt und über HTTP versendet.

- *Bibliotheken* - Um Entwicklern einen Teil der Arbeit für die Erstellung der Anfrage abzunehmen, existieren für die meisten Programmiersprachen Bibliotheken, sogenannte Libraries. Hier übersendet der Entwickler den durch die Bibliothek bereitgestellten Methoden die benötigten Informationen und diese erstellt und versendet eine korrekte Anfrage. Der Vorteil dabei ist, dass die Bibliotheksfunktionen zum einen intuitiver zu verwenden und zum anderen unabhängiger von Veränderungen auf Seiten der Cloud-Dienstanbieter sind. Ändert sich der Methodenaufruf auf Seiten des Cloud-Anbieters wird lediglich die Bibliothek geändert, nicht aber der Methodenaufruf des Programms an die Bibliothek.

Innerhalb der Bibliotheken stellen die *Multi Cloud Bibliotheken*, die sogenannten Multi Cloud Libraries (MCLs), eine Besonderheit dar. Insbesondere fehlende API

- **Abstraktionsschicht** - Hier wird die Generierung der Anfrage durch Einfügen einer Abstraktionsschicht komplett aus der Programmierenebene ausgelagert und so das Verwenden vollkommen generischer Anfragen ermöglicht. Im Gegensatz zu den MCLs, die einen eingeschränkten Funktionsumfang mit Bezug zu einer bestimmten Ressourcenart (Rechen, Speicher, etc.) aufweisen, unterstützt die Abstraktionsschicht ein viel breiteres funktionales Spektrum und den generischen Aufruf von komplexer aufgebauten Cloud-Diensten. Zudem kann eine Abstraktionsschicht auch von einem anderen Cloud-Anbieter in Form eines Cloud Broker zur Integration angeboten werden. Des Weiteren ist die Abstraktionsschicht - im Vergleich mit einer MCL - unabhängig von einer konkreten Programmiersprache.

Das Diagramm zeigt die Interaktion zwischen Cloud Diensten und Entwicklern. Oben befindet sich ein Kasten für den **Cloud Dienst Kunde**. Drei Entwickler (A, B, C) sind dargestellt, die jeweils ein Programm (A, B, C) entwickeln. Entwickler A nutzt eine **Web GUI**, während B und C eine **Bibliothek** und eine **Abstraktions-schicht** (mit **Treiber**) verwenden. Alle Entwickler kommunizieren über eine **Remote Desktop**-Verbindung mit einer Cloud, die die **Anbieter X, Y und Z** enthält. Eine **Legende** rechts unten erklärt die Verbindungstypen: gestrichelte Pfeile für generische Anfragen, durchgezogene Pfeile für API-spezifische Anfragen und durchgezogene Pfeile für sonstige Anfragen.

51

4 Cloud Computing

Im Folgenden werden die generischen Zugriffsarten, die im vorherigen Abschnitt definiert wurden, auf Ihre Anwendung und Umsetzung im Bereich Cloud Computing untersucht.

Hier sind insbesondere folgende Fragen von speziellem Interesse:

- *REST oder SOAP* - RESTful APIs haben hier das größte Potenzial. Im Vergleich mit SOAP besteht ein Vorteil darin, dass einige RESTful APIs die Möglichkeit bieten, Informationen in von XML verschiedenen Formaten zu versenden, was das eventuell nötige Übersetzen der Nachrichten überflüssig macht.
- *API oder Bibliothek* - Die Bibliotheken sind sehr simpel in der Verwendung, generieren sie doch durch einfache Methodenaufrufe selbständig korrekte Anfragen. Eine direkte API Anfrage ist etwas aufwändiger zu erstellen, da auf Details geachtet werden muss. In der Praxis hängt die Entscheidung zu allererst von der Verfügbarkeit einer gesuchten Bibliothek ab.
- *Wahl der Programmiersprache* - Hier steht Java an erster Stelle. Wenn es für einen Cloud Computing-Dienst eine Bibliothek gibt, ist eine für Java in den meisten Fällen vertreten. Dazu kommen die Multi Cloud Bibliotheken jClouds und Dasein sowie die Java Version von libcloud. Für Python, C# bzw. .NET, Ruby, PHP und Perl bieten ebenfalls einige Anbieter Bibliotheken an.

			Zugriffsarten			Formate			Bibliotheken						Multi Cloud Bibliotheken							
			Virtualisierung	Web GUI	Terminal-Sitzung	API	XML	JSON	C/S	text/plain	Java	Python	Ruby	C#/.NET	PHP	Perl	erlang	Delcloud	jclouds [Java]	Dasein [Java]	libcloud [Java]	libcloud [Python]
Anbieter	Produktname																					
Cloud Computing	Amazon	Elastic Compute Cloud	Xen	x	x	SOAP/Query	x				x	x	x	x	x	x	x	x	x	x	x	x
	Cloud Central	Servers	Xen	x		REST					x	x	x	x	x							
	ElasticHosts		KVM	x	x	REST		x		x												
	Flexiant	FlexiScale	Xen	x	x	SOAP/Query	x				x	x	x	x	x	x						
	GoGrid	Cloud Servers	Xen	x	x	Query	x	x	x		x	x	x	x	x	x			x	x		x
	IBM	Smart Business (Dev. and Trust)	RHEV	x	x	REST					x										x	x
	Rackspace	Cloud Servers	Xen	x	x	REST	x	x			x								x	x	x	x
	ReliaCloud	Cloud Servers		x		Query	x	x													x	
	Verizon	Computing as a Service	VMWare HA	x	x	REST	x															
	voxel	VoxCLOUD	Xen	x		REST	x	x			x	x	x		x	x						x
VPS.NET	Daily Nodes	Xen			REST		x														x	
Software	Cloud.com	CloudStack																		x		
	Enmaly	Elastic Computing Platform																				x
	Eucalyptus																		x	x	x	x
	Nimbula	Direktor																		x		
	OpenNebula																					x
	Red Hat	Enterprise Virtualization Manager																	x			
	VMWare	vCloud [Express]																	x	x		x
	VMWare	vSphere																			x	
verschiedene VPS-Dienste																			x	x	x	

Abbildung 2: Zugriffsarten und verfügbare Bibliotheken für Cloud Computing-Dienste

Neben den generischen Zugriffsarten sind in der Praxis sind die vorkonfigurierten Betriebssysteme, die der jeweilige Anbieter bereitstellt, sowie der physische Standort ebenfalls wichtig (siehe Abbildung 3).

		Betriebssysteme																		Regionen					
		Red Hat Enterprise Linux	Ubuntu	CentOS	SUSE Linux Enterprise	CloudLinux	Debian	Fedora	Knoptix	FreeBSD	OpenSolaris	Arch	Gentoo	ClearOS	Oracle Enterprise Linux	MS SQL Server 2005	MS SQL Server 2008	MS Windows Server 2003	MS Windows Server 2008	US (Ost)	US (Zentral)	US (West)	Europa	Australien	Silberstein
Anbieter	Produktname																								
Cloud Computing	Amazon Elastic Compute Cloud	x	x	x	x		x	x			x	x	x	x	x	x	x	x	x	x		x	x		
	Cloud Central Servers		x	x			x										x		x						
	ElasticHosts		x	x			x	x	x	x	x						x	x			x	x			
	Flexiant FlexiScale		x	x			x										x	x							
	GoGrid Cloud Servers	x	x	x											x	x	x	x				x			
	IBM Smart Business (Dev. and Trust)	x			x											x	x	x	x			x			
	Rackspace Cloud Servers	x	x	x			x	x				x	x	x	x		x	x	x	x					
	ReliaCloud Cloud Servers			x	x		x			x								x	x		x				
	Verizon Computing as a Service	x	x	x														x	x						
	voxel VoxCLOUD		x	x	x		x	x													x		x		x
VPS.NET Daily Nodes		x	x		x	x						x	x				x	x							

Abbildung 3: Verfügbare Betriebssysteme und Regionen von Cloud Computing-Diensten

5 Cloud-Speicher

Für die vorgestellten Zugriffsarten wird in diesem Abschnitt untersucht, ob und wie weit diese für vorhandene Cloud-Speicher Dienste umgesetzt wurden. Insbesondere folgende Fragen sind interessant:

- **BLOB oder Block** - Aufgrund der Flexibilität ist BLOB-Speicher für die meisten Aufgaben besser geeignet als Blockspeicher. Viele Vorteile von Blockspeicher (Einbindbarkeit/Ordnerhierarchie) können durch Tools auf BLOB-Speicher übertragen werden.
- **REST oder SOAP** - Im Vergleich der Übertragungskonzepte ist REST klar zu bevorzugen. „85 Prozent des Marktes wendet sich REST zu während SOAP verschwindet“ [Wen09], da RESTful APIs einfacher zu verwenden sind als SOAP APIs, ohne dabei wesentlich an Funktionalität einzubüßen.
- **API oder Bibliothek** - Die Abwägung zwischen den beiden Zugriffsarten ist im Grunde ähnlich wie bei Cloud Computing-Diensten, da ebenfalls nicht für alle Dienst-Programmiersprache-Kombinationen passende Bibliotheken verfügbar sind.
- **Wahl der Programmiersprache** - Gegenwärtig gibt es für Java die meiste Unterstützung, sei es in Form von Bibliotheken für einzelne Dienste oder Multi Cloud Bibliotheken. Mit Dasein, CloudLoop und jClouds basieren drei Multi Cloud Bibliotheken auf Java. Daneben ist PHP die Sprache, die sich durch ihre Verwendung in der simplecloud Bibliothek von den Sprachen C# bzw. .NET, Ruby und Python absetzt.

Abbildung 4 stellt die wichtigsten Eigenschaften der untersuchten Cloud-Speicheranbieter gegenüber.

			Zugriffsarten			Formate		Bibliotheken						Multi Cloud Bibliotheken								
			Web GUI	Terminal-Sitzung	API	XML	JSON	text/plain	Java	Python	Ruby	C#/.NET	PHP	JavaScript	Objective C	simplecloud [PHP]	iclouds [Java]	Dasein [Java]	CloudLoop [Java]	SMEStorage [REST]		
BLOB-Speicher	Anbieter	Produktname	x	x	REST/SOAP	x			x		x	x	x			x	x	x	x	x		
	Amazon	Simple Storage Service	x	x	REST	x				x							x					
	Google	Storage for Developers		x						x												
	Mezeo	Cloud Storage Platform		x	REST	x	x		x	x					x					x		
	Microsoft	Windows Azure Storage		x	REST	x			x				x	x			x	x	x			
	Nirvanix	Storage Delivery Network		x	REST/SOAP	x	x	x	x	x	x	x	x	x		x			x			
	Rackspace	Cloud Files	x	x	REST	x	x	x	x	x	x	x	x			x	x	x	x	x		
Software	weitere BLOB-Speicherdienste																x	x				
	EMC	Atmos																x	x	x		
	Eucalyptus	Walrus																x				
	Oracle	Sun Cloud (eingestellt)																		x		

Abbildung 4: Zugriffsarten und Bibliotheken für Cloud-Speicheranbieter

6 Platform und Software as a Service

Im Vergleich zu Cloud Computing und Cloud-Speicher ist eine Differenzierung bei PaaS und SaaS schwieriger. Grund dafür ist zum einen, dass beide Ebenen darauf ausgelegt sind dem Benutzer Zugriffsarten, Ressourcenverwaltung und Konfiguration weitestgehend vorzuenthalten und zum anderen, dass hier individuelle Produkte einzelner Anbieter vertreten sind, die es nur begrenzt ermöglichen allgemeine Konzepte zu finden.

Von besonderem Interesse ist hier die Wahl der Programmiersprache. PaaS-Dienste können zwar prinzipiell ein breites Spektrum von Applikationen hosten. Die praktische Untersuchung (Abbildung 5) jedoch hat gezeigt, dass sehr wohl eine gewisse Sprachabhängigkeit besteht. Es ist zu erkennen, dass Java sowie Ruby überwiegend zu empfehlen sind. Entwicklerwerkzeuge für Python, PHP und für das .NET-Framework sind jedoch ebenfalls verfügbar.

		Zugriffsarten			Verfügbare Sprachen				
		Web GUI	Terminal-Sitzung	Remote Desktop	Java	Python	Ruby	C#/.NET	PHP
PaaS	Microsoft Windows Azure Storage	x		x	x		x	x	x
	Google Google App Engine	x			x	x			
	Salesforce Cloud Storage Platform	x			x				
	Engine Yard Windows Azure Storage	x					x		
	Heroku Storage Delivery Network	x	x				x		

Abbildung 5: Übersicht einiger PaaS-Dienste

7 Fazit

Das Anfangs vorgestellte erweiterte SPI-Ebenen-Modell spaltete die einfache Aufteilung von Cloud-Diensten in IaaS, PaaS und SaaS weiter auf, indem IaaS in Computing und Speicher aufgeteilt wurde. Letzterer wurde dabei zusätzlich in die verschiedenen Konzepte BLOB- und Blockspeicher unterteilt. Hier war festzustellen, dass BLOB-Speicher empfehlenswerter sind, da dort die Konzepte der Cloud besser umgesetzt werden.

Als Schwerpunkt wurden Zugriffsmöglichkeiten auf Dienste und ihre Umsetzungen untersucht. Bei der Untersuchung von IaaS-Diensten konnte festgestellt werden, dass neben dem Zugriff via Web GUI pro Anbieter zum Teil sogar mehrere APIs und Bibliotheken für verschiedene Programmiersprachen verfügbar sind. Durch Untersuchen von RESTful und SOAP-basierten API Anfragen konnte festgestellt werden, dass sich REST-konforme Anfragen besser für die Kommunikation mit Cloud-Diensten eignen, da sie tendenziell simpler zu verfassen und zu verarbeiten sind. (Multi Cloud) Bibliotheken sind weit verbreitet und unterstützen in erster Linie Java-Programmierer. Daneben ist PHP die Sprache, die sich durch ihre Verwendung in der simplecloud Bibliothek von den Sprachen C# bzw. .NET, Ruby und Python absetzt. Mit PaaS- und SaaS-Ebene endete die Untersuchung von Implementierungen und Zugriffsarten von und auf Cloud-Dienste, wobei hier aufgrund der Verschiedenartigkeit der Dienste nur begrenzt allgemeine Aussagen getroffen werden können.

In Bezug auf Standards bleibt festzuhalten, dass es momentan für die jeweilige SPI-Ebene keinen Standard gibt, der von allen Cloud-Anbietern umgesetzt wird. Dies stellt eines der größten Probleme dar. An vielen Stellen der Kommunikation zwischen Kunden und Anbietern gibt es Hilfsmittel oder Abstraktionsebenen, die diese fehlenden Konventionen kaschieren. Auch wenn bereits ein Annähern der Dienste aneinander zu erkennen ist und die Entwicklung von Standards langsam nennenswerte Ergebnisse hervorbringt [CDM10] [MKL09] ist davon auszugehen, dass es noch einige Jahre dauern wird, bis die Entwicklung und das Anerkennen von Standards abgeschlossen sind.

8 Danksagung

Besonderen Dank gilt der Professur für Kommunikationssysteme und Internet-Dienste der Universität der Bundeswehr München, allen voran Frau Prof. Dr. Gabi Dreö Rodosek sowie dem Munich Network Management Team für die hilfreichen Diskussionen und wertvollen Kommentare bei der Erstellung dieser Publikation. Darüber hinaus möchten wir uns bei der NATO SC/4 SMI AHWG und der Europäischen Verteidigungsagentur für die Unterstützung und Mitwirkung bedanken. Unsere Forschung in diesem Bereich wurde auch teilweise in Kooperation mit dem Bundesamt für Sicherheit in der Informationstechnik durchgeführt.

Literatur

- [AZS11] Windows Azure Storage — Windows Azure Platform. Online-Quelle, 2011. <https://www.microsoft.com/windowsazure/storage/default.aspx>.
- [Bia08] Randy Bias. Virtual, Cloud, Datacenters? Online-Quelle, April 2008. <http://cloudscaling.com/blog/cloud-computing/virtual-cloud-datacenters>.
- [CDM10] Storage Networking Industry Association. Online-Quelle, 2010. <http://cdmi.sniacloud.com/>.
- [Chi07] D. Chisnall. *The definitive guide to the xen hypervisor*. Prentice Hall Press Upper Saddle River, NJ, USA, 2007.
- [GGW09] Scaling Your Internet Business. Online-Quelle, März 2009. <http://www.gogrid.com/downloads/GoGrid-scaling-your-internet-business.pdf>.
- [Hes10] Kenneth Hess. Top 10 Virtualization Technology Companies. Online-Quelle, April 2010. <http://www.serverwatch.com/trends/article.php/3877576/Top-10-Virtualization-Technology-Companies.htm>.
- [jcl10] jclouds. Online-Quelle, 2010. <http://jclouds.org/>.
- [MKL09] T. Mather, S. Kumaraswamy und S. Latif. *Cloud security and privacy: An enterprise perspective on risks and compliance*. O'Reilly Media, Inc., 2009.
- [S310] Amazon Simple Storage Service (Amazon S3). Online-Quelle, 2010. <http://aws.amazon.com/s3>.
- [SfD11] Google Storage for Developers. Online-Quelle, 2011. <https://code.google.com/apis/storage>.
- [SFS11] Sales Cloud 2. Online-Quelle, 2011. <http://www.salesforce.com/crm/sales-force-automation/>.
- [SFW10] Force.com: Build and run business apps 5 times faster at half the cost. Online-Quelle, 2010. http://www.salesforce.com/assets/pdf/datasheets/DS_Forcedotcom.pdf.
- [TSW] White Paper Cloud Computing. Online-Quelle. http://download.sczm.t-systems.de/t-systems.ch/de/StaticPage/62/79/00/627900_White-Paper-Cloud-Computing.
- [War08] Keith Ward. More Azure Hypervisor Details. Online-Quelle, November 2008. <http://virtualizationreview.com/blogs/mental-ward/2008/11/more-azure-hypervisor-details.aspx>.
- [Wen09] Jerome Wendt. Archiving data to cloud storage: How to choose the right cloud storage provider. Online-Quelle, August 2009. http://searchstorage.techtarget.com/tip/0,289483,sid5_gci1365069_mem1,00.html.
- [Xig10] Getting Started - Market Data Feed - Financial Web Service - On-Demand - Xignite. Online-Quelle, 2010. <http://www.xignite.com/Support/GettingStarted.aspx>.
- [Zet10] Zetta Enterprise Network Storage Service. Online-Quelle, 2010. <http://www.zetta.net/zettaStorageService.php>.

Virtuelle Forschungsumgebungen

Toolset zur Planung und Qualitätssicherung von verteilten virtuellen Netzwerkstrukturen

Martin Gründl, Susanne Naegele-Jackson

Friedrich-Alexander Universität Erlangen-Nürnberg

Martensstrasse 1, 91058 Erlangen

{martin.gruendl, susanne.naegele-jackson}@rrze.uni-erlangen.de

Zusammenfassung: Durch die zunehmende Verbreitung von leicht handhabbarer Software für anspruchsvolle Anwendungen wie Videoconferencing oder Server-Virtualisierung sind neue verteilte virtuelle Netzwerkstrukturen und zwischenbetriebliche Vernetzungen möglich. Bei solchen verteilten virtuellen Strukturen ist es für den Benutzer von großem Vorteil, wenn schon im Vorfeld abgeklärt werden kann, ob gewisse Applikationen über alle verteilten Standpunkte nutzbar sind und auch in der entsprechenden Dienstgüte gewährleistet werden können. Diese Arbeit beschreibt ein Toolset, welches dazu verwendet werden kann, Schwachstellen in einem Multi-Domain-Netz aufzuzeigen, Dienstgüteparameter zu messen und die Erfüllung der Anforderungen einer Applikation für eine verteilte virtuelle Organisationsstruktur zu testen. Weiterhin werden Mechanismen zur Überprüfung von Protokollparametern und deren prinzipieller Funktion sowie Werkzeuge zur Bewertung der IPv6-Konnektivität angeboten. Um den verteilten Einsatz des Systems zu erleichtern, liegt ein besonderer Fokus auf der Integration verschiedener Sensor-Implementierungen, um sowohl kostengünstige Messpunkte mit Hilfe von Plug Computern oder virtuellen Maschinen als auch hochpräzise Sensoren für hohe Bandbreiten-Anforderungen realisieren zu können.

1 Einleitung

Die informationstechnische Vernetzung bietet Firmen die Möglichkeit, ihre Betriebsprozesse über neue virtuelle Organisationsstrukturen zu verteilen. Bei einer solchen Verteilung kann es sich beispielsweise um eine Teilmenge von Firmenstandorten handeln, die alle an einem bestimmten Projekt arbeiten und zusätzlich auch externe Berater in diese projektabhängige Organisationsstruktur einbinden möchten [We06]. Das Netz als Verbindungsglied dieser Struktur stellt die Basis der gemeinsamen Projektarbeit dar und nimmt dadurch einen sehr hohen Stellenwert ein. Störungen im Netz oder ungenügende Dienstqualität für einzelne Standorte der virtuellen Netzwerkstruktur können zu einer massiven Beeinträchtigung der Arbeitsabläufe innerhalb des gesamten Projektes führen. Es ist daher von besonderem Vorteil, wenn Netzwerkzeuge zur Verfügung stehen, die zum einen die Fehlersuche bei Störungen im Netz vereinfachen, auf der anderen Seite aber auch Möglichkeiten bieten, Applikationen schon im Vorfeld auf ihre Einsetzbarkeit an den einzelnen Standorten zu testen. Ein

typischer Fall hierfür sind Videokonferenzen, die über alle Standorte hinweg in ausreichend guter Dienstqualität durchführbar sein sollten. Der Benutzerkreis möchte sich vor der Nutzung der Applikation einen Überblick darüber verschaffen, ob über die Teilbereiche der Netzinfrastruktur genügend Ressourcen zur Verfügung stehen, so dass die erforderliche Dienstgüte der Anwendung gewährleistet werden kann.

Diese Arbeit beschreibt TSS, einen Netzwerk Troubleshooting Service, der Funktionstests ermöglicht, welche für den Einsatz in verteilten virtuellen Strukturen von Interesse sind. Zum Funktionsumfang von TSS zählen unter anderem prinzipielle Netzfunktionstests wie `ping` und `traceroute`, darüber hinaus kann aber mit TSS auch die prinzipielle Funktionalität von Netzservices wie DHCP oder DNS überprüft werden. TSS erlaubt außerdem die Bestimmung von protokoll- und anwendungsspezifischen Parametern und bietet die Möglichkeit, IP Performance Metrics zur Qualitätssicherung einzusetzen.

2 Verwandte Arbeiten

Bei TSS handelt es sich um eine Weiterentwicklung von perfSONAR-Lite Troubleshooting Service (TSS) [NGH10], das im Rahmen des EGEE-III (Enabling Grids for E-science) Projektes entwickelt wurde. Da EGEE als eines der führenden Grid-Projekte mehr als 280 Einrichtungen in Europa vernetzte, war es notwendig, eine Lösung zur Fehlersuche auf Netzebene zur Verfügung zu haben, um schnell bei akuten Problemen eine leicht handhabbare Untersuchungssoftware einsetzen zu können. Diese am Regionalen Rechenzentrum Erlangen entwickelte Software basierte auf Teilen des im GN2/GN3-Projekt entwickelten Monitoringsystems perfSONAR [BBD05,HKM08] und stellte Tools wie `ping`, `traceroute`, `dns lookup`, `nmap` [Nm11] sowie `bwctl` [Bw11] für Bandbreitenmessungen zur Verfügung. In dieser Arbeit wird nun die vom DFN-Verein finanzierte Weiterentwicklung von TSS vorgestellt, die außer den bereits vorhandenen für EGEE-III entwickelten Tools weitere Werkzeuge enthält, die nicht nur für eine Fehlerdiagnose entscheidend sind, sondern darüber hinaus zusätzliche Netzfunktionalitäten prüfen können und daher vor allem auch bei der Planung von verteilten virtuellen Organisationsstrukturen oder zur Überprüfung von Qualitätsanforderungen bei Anwendungen hilfreich sind.

Ein vergleichbares Schnelldiagnose-System findet man in der Literatur unter NDT (Network Diagnostic Tool [Nd11]): NDT bietet unter anderem Tests, um Ethernet-Duplex-Einstellungen zu überprüfen oder Links zu ermitteln, die nicht den zugesagten Datendurchsatz erreichen. Allerdings führt NDT diese Tests immer zwischen dem Desktop Computer des Benutzers und dem NDT-Server durch. TSS hingegen erlaubt die Schnelldiagnose losgelöst vom PC eines Benutzers zwischen zwei beliebigen Sensoren im Netz. Dadurch wird eine Schnelldiagnose im Fehlerfall flexibler und auch Qualitätsanforderungen für Anwendungen lassen sich präzise über den betroffenen Pfad bestimmen.

Im folgenden Abschnitt wird zunächst die Implementierung von TSS erläutert, während Abschnitt 4 die möglichen Anwendungsszenarien der TSS Werkzeuge im De-

tails beschreibt. Schließlich werden in Abschnitt 5 die Kernaussagen zusammengefasst und ein Ausblick auf mögliche weitere Entwicklungen gegeben.

3 Architektur

Die grundlegende Architektur des TSS-Systems setzt sich aus einem Server und den für die eigentlichen Messungen genutzten Sensoren zusammen.

3.1 Zentraler Server

Der Webserver als zentrale Komponente kommuniziert über einen abgesicherten Kanal mit den Sensoren und stellt die Schnittstelle für die Benutzer zur Verfügung. Der Ablauf einer entsprechenden Messung ist in Abbildung 1 dargestellt.

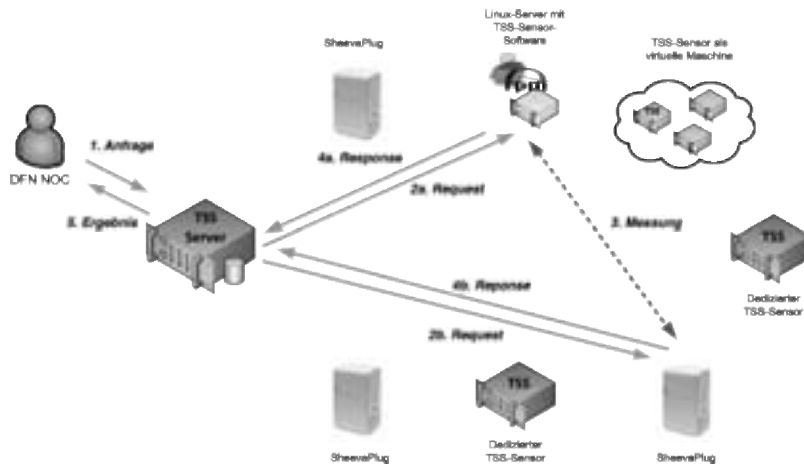


Abbildung 1: Ablauf einer Messung

Das Webserver Interface führt eine zugangsberechtigte Person durch die Auswahl der entsprechenden Messtypen und Parameter (1) und leitet schließlich die entsprechenden Mess-Aktionen über einen SSL-Kanal an die Sensoren des Systems weiter (2). Diese wiederum nehmen ausschließlich über diesen Kanal Befehle entgegen, führen sie aus (3) und liefern die Ergebnisse wieder an den Server zurück (4). Schließlich werden die Resultate der Messung dem Benutzer präsentiert (5).

3.2 Sensoren

Um einen möglichst breiten Bereich an Anwendungsfällen abdecken zu können, ist es ein zentrales Ziel des hier vorgestellten Systems, an verschiedene Einsatzszenarien angepasste Sensoren zur Verfügung zu stellen (Abb. 2). Grundsätzlich handelt es sich

bei einem solchen Sensor um einen speziellen Client zur Durchführung von Messungen und Funktionstests zwischen zwei Messpunkten; die dazugehörigen Ergebnisse werden dann über das Webserver Interface dem Anwender übermittelt.

Neben der für das EGEE-III Projekt entwickelten Software-Lösung, welche auf einem bestehenden Linux-Server in wenigen Schritten installiert werden kann und daher auch keine zentrale Administration des Systems erlaubt, sind für mit begrenzten Mitteln ausgestattete Vorhaben oder noch in der Planungsphase befindliche Projekte vor allem die Sensoren auf Basis der SheevaPlug-Familie [Sh11] interessant. Dieses Gerät, welches trotz seiner Kompaktheit und der sehr geringen Anschaffungskosten ein vollständiges Linux-System zur Verfügung stellt, benötigt lediglich eine Netzwerkverbindung sowie einen Stromanschluss, um seinen Betrieb aufnehmen zu können.

Da das Betriebssystem dieser Sensoren von einer SD-Karte gestartet wird, ist im Fehlerfall keine langwierige und kostenintensive Wiederherstellung erforderlich. Im für den Benutzer einfachsten Fall wird lediglich die SD-Karte getauscht, alternativ kann sie natürlich auch mit Hilfe eines für registrierte Nutzer zugänglichen Images auf einen definierten Zustand zurückgesetzt werden.

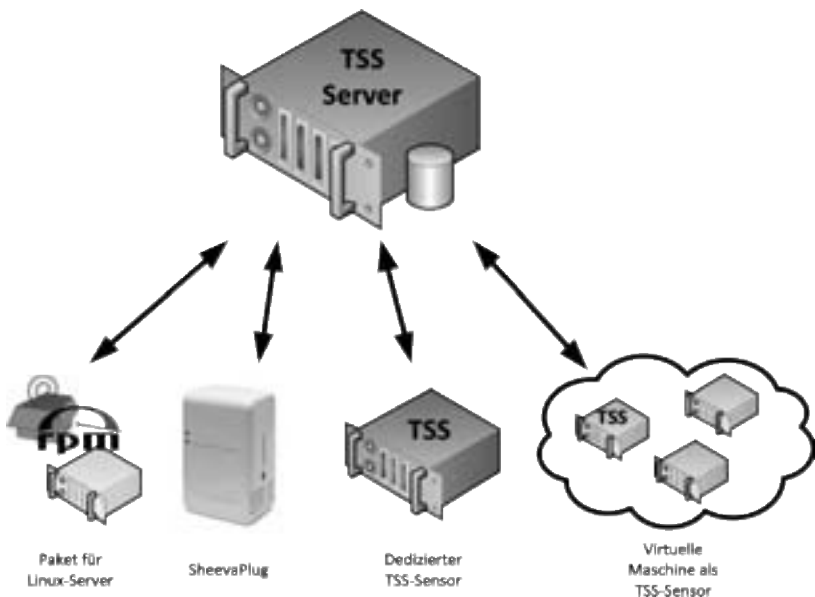


Abbildung 2: Verfügbare Sensoren

Natürlich besteht weiterhin die Möglichkeit, einen vollwertigen Server, ähnlich zu den erfolgreich in nationalen und internationalen Forschungsnetzen betriebenen HADES-Messrechnern [Ha11], in das System zu integrieren. Ein solcher HADES Messrechner erlaubt aufgrund seiner Bauform und seiner erweiterten technischen

Ausstattung vor allem im Bereich des Zeitverhaltens mit Hilfe einer GPS-Karte sehr viel präzisere Messungen und kann beispielsweise auch - ausgerüstet mit einer entsprechend leistungsfähigen Netzwerkkarte - Durchsatzmessungen bis in den Bereich von mehreren zehn Gigabit pro Sekunde durchführen.

Gerade im Rahmen der rasant zunehmenden Virtualisierung ist es darüber hinaus möglich, eine mit der entsprechenden Software ausgestattete virtuelle Maschine ebenfalls als Sensor zu integrieren, um auch innerhalb derartiger virtueller Umgebungen die diversen Möglichkeiten des TSS-Systems nutzen zu können.

3.3 Verfügbare Werkzeuge

Um die zahlreichen Anwendungsfälle des Systems abdecken zu können, wurde eine breit gefächerte Auswahl an Werkzeugen integriert, welche in der folgenden Tabelle 1, jeweils nach Funktionalität gegliedert, dargestellt ist.

Tabelle 1: Werkzeuge

Konnektivität	IPv6-Unterstützung	Protokoll- und anwendungs-spezifisch	IPPM-Metriken und Qualitäts-Indikatoren
ping	ping6	tracpath	OWAMP
tracroute	tracroute6	tracpath6	iperf
host / nslookup	Router Discovery	netcat	
DHCP-Diagnose		nmap	
NTP			

3.4 Verwaltung der Berechtigungen

Die Authentifizierung der Benutzer erfolgt über X.509-Zertifikate, um eine eindeutige und zuverlässige Identifikation sicherstellen zu können. Die Berechtigungen innerhalb des Systems können wiederum spezifisch auf Ebene der Sensoren zugewiesen werden. Einem nicht privilegierten Benutzer des Systems ist es somit möglich, jeweils Messungen ausgehend von einem für ihn freigegebenen Sensor zu allen übrigen Sensoren auszulösen, sowie in entgegengesetzter Richtung von allen übrigen Sensoren zu den für ihn freigegebenen Sensoren durchzuführen.

4 Anwendungsszenarien

Um die Arbeit der mit dem Aufbau und dem Betrieb eines Netzwerkes befassten Personen zu unterstützen, werden mit Hilfe der bereits beschriebenen Architektur zahlreiche Parameter, welche einen Netzwerkanschluss charakterisieren, ermittelt, um

sowohl bei der Planung eines Netzes als auch bei der späteren Behebung von Anomalien oder Fehlern möglichst effizient Hinweise auf Fehlerquellen zu erhalten. Im Folgenden werden einige Szenarien für den Einsatz des TSS-Systems skizziert.

4.1 Überprüfung der prinzipiellen Funktionsfähigkeit und Qualität der Netzanbindung

Im Rahmen der Planung neuer PoPs oder neuer Standorte von virtuellen Organisationsstrukturen ist zunächst eine uneingeschränkte Konnektivität der dortigen Netzanbindungen abzuklären. Neben den trivialen Mechanismen zur Überprüfung der Erreichbarkeit wie ping ist es hier beispielsweise von Interesse, ob alle bestehenden oder geplanten PoPs von dem jeweils untersuchten Standort aus erreicht werden können oder ob hier eventuell gewollte oder durch Störungen verursachte Einschränkungen vorliegen. Neben der prinzipiellen Erreichbarkeit der Gegenstellen ist natürlich auch die Nutzung gleichbleibender Wege innerhalb der Transit-Netzwerke erstrebenswert. Auch wenn dies prinzipiell keine Anforderung im Rahmen des IPv4-Standards ist, da dieser paketvermittelt arbeitet und somit explizit auf die variable Wegführung der IP-Pakete ausgelegt ist, sind aus Gründen des Zeitverhaltens und anderer potentiell störender Nebeneffekte konstant bleibende Routen, welche sich lediglich im Fehlerfall oder während Wartungsarbeiten verändern, ein relevantes Qualitätsmerkmal. Zu diesem Zweck kommen, abhängig von den konkreten Anforderungen und zu bestimmenden Parametern, Vertreter der unter dem Oberbegriff traceroute zusammengefassten Werkzeuge zum Einsatz.

Darüber hinaus ist natürlich auch die einwandfreie Funktionalität von Netzdiensten wie dem Domain Name System DNS und dem Dynamic Host Configuration Protocol DHCP (Abb. 3) eine Voraussetzung für eine zuverlässige und leistungsfähige Netzanbindung. Daher können hier ebenfalls einige grundlegende Funktionen geprüft werden, beispielsweise wie sich der lokal angebotene DNS-Server im Falle wiederholter identischer Anfragen verhält, um eventuelle Caching-Effekte erkennen und gegebenenfalls berücksichtigen zu können.



Abbildung 3: Screenshot DHCP-Diagnose

4.2 Qualität der IPv6-Konnektivität

Die bisher genannten Messungen bekommen, vor allem aufgrund der Erschöpfung des verfügbaren IPv4-Adressraumes, mit der rasch zunehmenden Verbreitung des Protokolls IPv6 eine neue Dimension, da in naher Zukunft sehr viele Einrichtungen ihre entsprechenden Netzstrukturen an das prinzipiell bereits seit mehr als 10 Jahren standardisierte und verfügbare Protokoll anpassen und den zuverlässigen Betrieb dieses Protokolls sicherstellen müssen. Da dieser Vorgang bisher keineswegs abgeschlossen ist, werden momentan noch diverse Mechanismen genutzt, um auch ohne native IPv6-Konnektivität bereits das entsprechende Protokoll nutzen zu können. Diese Technologien weisen durch ihre Aufgabe natürlich einen inhärenten Einfluss auf das Verhalten der darüber verwendeten Protokolle und im Speziellen auf das Zeitverhalten auf, welches beispielsweise für zeitkritische Projekte durchaus relevante Auswirkungen haben kann.

Daher ist an dieser Stelle ein wesentliches Qualitätsmerkmal der IPv6-Konnektivität, ob diese uneingeschränkt von der den Netzanschluss zur Verfügung stellenden Infrastruktur unterstützt wird oder vorerst der Einsatz der beschriebenen Mechanismen nötig ist. Durch die Überprüfung diverser Parameter am topologischen Standort des Senders und einige Messungen kann sowohl das prinzipielle Verfahren der IPv6-Anbindung als auch deren Qualität gemessen und überprüft werden.

4.3 Bestimmung von protokoll- und anwendungsspezifischen Parametern

Neben der prinzipiellen Funktionalität einer Netzanbindung, welche oben beschrieben wurde, sind weitere Parameter für die optimale Nutzung von Protokollen höherer Schichten und Anwendungen relevant. Um beispielsweise eine möglichst effiziente Übertragung von großen Datenmengen realisieren zu können, ist der Parameter Maximum Transmission Unit (MTU) von großer Bedeutung, da dieser Parameter die maximale Größe eines IP-Paketes beschreibt. Im optimalen Fall ist die Nutzung von Jumbo Frames mit einer Größe von bis zu 9000 Bytes möglich, wodurch sich sowohl volumenmäßig als auch auf technischer Ebene der Overhead gravierend reduzieren lässt. Daher sind Werkzeuge zur Bestimmung der MTU auf einem Pfad ebenfalls Bestandteil der hier beschriebenen Lösung.

Mit Hilfe des hier vorgestellten Werkzeugs kann somit eine Verbindung mit den entsprechenden Quell- und Zielpoints beziehungsweise den Protokollnummern simuliert werden, um einen prinzipiellen Verbindungsaufbau nachstellen zu können. Neben der Funktionsfähigkeit von Anwendungen, welche im späteren Betrieb eingesetzt werden sollen, und eventuell störenden Netzkomponenten wie beispielsweise Access-Listen oder Firewalls wird mit diesem Mechanismus zudem überprüft, ob die für die Administration der TSS-Sensoren nötigen Verbindungen sowie die für Messungen benötigten Werkzeuge und ihre jeweiligen Protokolle über den entsprechenden Netzanschluss nutzbar sind.

4.4 Bestimmung von IPPM-Metriken und weiterer Qualitätsindikatoren

Die Qualität einer Netzanbindung wird neben den für eine grundlegende Funktion nötigen Protokollen und Parametern im Wesentlichen durch Metriken, welche größtenteils unter dem Oberbegriff IP Performance Metrics zusammengefasst werden, erfasst und bewertet. Vor allem für die Übertragung von Multimediadaten zur Echtzeitkommunikation sowie für zeitkritische Steuerungsaufgaben haben die Parameter One-Way Delay und One-Way Delay Variation erhebliche Bedeutung, da sie das zeitliche Verhalten einer Verbindung im Bezug auf eventuelle Verzögerungen oder Schwankungen im Zeitverhalten charakterisieren.

Im Rahmen des Einsatzes von Voice over IP ist dies beispielsweise relevant, um die Größen der für den Ausgleich von Laufzeitschwankungen nötigen Pufferspeicher vorab oder auch dynamisch während des Betriebs festlegen zu können. Ein positiver Nebeneffekt einer derartigen Messung ist weiterhin eine Erkennung von gegebenenfalls auftretenden Paketverlusten.

Zu diesem Zweck wird das One-Way Active Measurement Protocol OWAMP eingesetzt (Abb. 4), welches beide Parameter auf der Grundlage entsprechender RFCs [STK06] bestimmt. Natürlich wäre prinzipiell auch der Einsatz des aktiven HADES-Messsystems denkbar, allerdings ist dies in seiner momentanen Form auf vorkonfigurierte periodische Messungen ausgerichtet und unterstützt keine Ad-hoc-Messungen.



Abbildung 4: Screenshot OWAMP-Messung

Neben dem Zeitverhalten des über die Internetverbindung fließenden Verkehrs ist natürlich gerade für wissenschaftliche Projekte und HPC-Anwendungen, wie beispielsweise die mehrstufige Datenverarbeitung der LHC-Experimente mit ihren enormen Datenmengen, der Durchsatz eines Netzanschlusses von hoher Wichtigkeit. Zur Bestimmung dieser Parameter kommt momentan das Werkzeug `iperf` [Ip11] zum Einsatz.

5 Fazit und Ausblick

Mit dem vorgestellten Werkzeug wurde eine Lösung geschaffen, welche eine Reihe von Anforderungen im Rahmen der Netzwerkplanung und der Erweiterung bestehender Netze sowie der Fehlerbehebung in entsprechenden Netzen abdeckt. Durch die verteilte Infrastruktur, welche durch ein zentrales System gesteuert wird und einen an die jeweiligen örtlichen Gegebenheiten und technischen Anforderungen angepassten Einsatz von Messpunkten verschiedener Komplexität ermöglicht, wird ein sehr breites Spektrum abgedeckt. Während bei sehr verteilten Projekten im Planungsstadium eine kostengünstige und portable Lösung mit Sicherheit vorgezogen wird, können für den Betrieb wichtiger Projekte mit hohen Qualitätsansprüchen bei Bedarf Messpunkte eingesetzt werden, welche deutlich komplexere und auch exaktere Messungen erlauben.

Um die Flexibilität und den Nutzwert des hier beschriebenen Systems weiter zu erhöhen, ist zukünftig die Integration weiterer Werkzeuge angedacht, wie beispielsweise eine anwendungsspezifische Simulation von Datenströmen mit geeigneten Werkzeugen.

Literatur

- [We06] Weiss, P., Virtuelle Organisationsstrukturen, Konzept, Analyse, Lösungen und Praxisbeispiele, Öffentlicher Abschluss-Workshop Projekt ARBEIT@VU, FZKI Forschungszentrum Informatik und Institut AIFB Universität Karlsruhe, 02.06.2006, http://www.virtuelleunternehmen.com/arbeit_at_vu/infothek/presentationen/20060602_workshop_abschluss.pdf
- [NGH10] Naegele-Jackson S., Gründl M., Hanemann A., PerfSONAR-Lite TSS: Schnelldiagnose von Netzverbindungen im EGEE-III-Projekt, 3. DFN-Forum „Verteilte Systeme im Wissenschaftsbereich“, 26. Und 27. Mai 2010, Universität Konstanz, Lecture Notes in Informatics (LNI) – Proceedings Series of the Gesellschaft für Informatik (GI), Paul Müller, Bernhard Neumair, Gabi Dreo Rodosek (Hrsg.), Volume P-166, Bonn 2010, pp. 127-136, ISBN 978-3-88579-260-4, <http://subs.emis.de/LNI/Proceedings/Proceedings166/P-166.pdf>
- [BBD05] Boote, J. W., Boyd, E. L., Durand, J., Hanemann, A., Kudarimoti, L., Lapacz, R., Swany, D. M., Zurawski, J., Trocha, S., PerfSONAR: A Service Oriented Architecture for Multi-Domain Network Monitoring, In *Proceedings of the Third International Conference on Service Oriented Computing*, 241–254, LNCS 3826, Springer Verlag, ACM Sigsoft, Sigweb, Amsterdam, The Netherlands, Dezember, 2005
- [HKM08] Hanemann, A., Kraft, S., Marcu, P., Reinwand, J., Reiser, H., Schmitz, D., Venus, V., *perfSONAR: Performance Monitoring in europäischen Forschungsnetzen*, In *Proceedings Erstes DFN-Forum Kommunikationstechnologien — Verteilte Systeme im Wissenschaftsbereich*, GI-Verlag/DFN, Kaiserslautern, Deutschland, Mai, 2008
- [Nm11] NMAP, <http://nmap.org/>, 06.04.2011
- [Bw11] Bandwidth Test Controller (BWCTL), <http://e2epi.internet2.edu/bwctl/>, 06.04.2011
- [Nd11] Network Diagnostic Tool (NDT), <http://www.internet2.edu/performance/ndt/>, 06.04.2011
- [Sh11] SheevaPlug, <http://www.plugcomputer.org/>, 06.04.2011
- [Ha11] HADES Active Delay Evaluation System, <http://www.win-labor.dfn.de/>, 06.04.2011
- [STK06] Shalunov, S., Teitelbaum, B., Karp, A., Boote, J., Zekauskas, M., RFC 4656 – A One-Way Active Measurement Protocol (OWAMP), IETF Network Working Group, 2006
- [Ip11] Iperf, <http://iperf.sourceforge.net>, 06.04.2011

Vom Studiolo zur virtuellen Forschungsumgebung

Andreas Brennecke, Gudrun Oevel, Alexander Strothmann

Universität Paderborn
Zentrum für Informations- und Medientechnologien (IMT)
Warburger Str. 100
33098 Paderborn

andreas.brennecke@uni-paderborn.de
gudrun.oevel@uni-paderborn.de
alexander.strothmann@uni-paderborn.de

Abstract: Für das Fach Kunst- und Architekturgeschichte wird eine virtuelle Forschungsumgebung entwickelt. Die Forschungsumgebung soll einerseits den spezifischen Forschungsbereich der Arbeit mit Bildern und multimedialen Dokumenten unterstützen, andererseits aber auch Funktionen für allgemeine, nicht fachspezifische Aspekte des Forschungsprozesses bieten. Letztere sollen so gestaltet werden, dass die implementierten Werkzeuge von vielen Fächern genutzt werden können. Basis der virtuellen Forschungsumgebung sind ein Open-Source-Framework zur Realisierung virtueller Wissensräume, die vorhandene Dienste-Infrastruktur sowie die kommerzielle „Geschäftskollaborationsplattform“ SharePoint.

1 Einleitung

Während mit dem Schlagwort „Virtuelle Forschungsumgebungen“ insbesondere fachspezifische oder an den Forschungsgegenstand angepasste Anwendungen bezeichnet werden, fassen wir diesen Begriff weiter. Nach unserem Verständnis gehört insbesondere auch die Organisation des Forschungsbetriebs dazu und sollte technologisch unterstützt werden. Virtuelle Forschungsumgebungen ermöglichen nicht nur zeit- und ortsunabhängiges Arbeiten, sondern stellen den Wissenschaftlern einen möglichst integrierten virtuellen Arbeitsplatz zur Verfügung.

Im Rahmen eines von der DFG geförderten Projekts [St09] wird an der Universität Paderborn eine virtuelle Forschungsumgebung aufgebaut, die diese Anforderungen realisieren soll. Das Zentrum für Informations- und Medientechnologien (IMT) der Universität Paderborn ist am Projekt maßgeblich beteiligt, nicht nur, um die nachhaltige Verankerung der Forschungsumgebung nach Ablauf des Projekts zu gewährleisten, sondern auch, um Know-how bei der Planung und Entwicklung von Anwendungen aufzubauen. Als zentrale Einrichtung für die Bereiche IT und Medien ist das IMT nicht nur ein Anbieter standardisierter Dienstleistungen, sondern wird zunehmend als Partner zur Realisierung komplexer angepasster Lösungen gefordert.

2 Vom Studiolo zum virtuellen Wissensraum

Ein Studiolo ist ein in der Renaissance entstandener spezieller Typ von Räumen, die mit Kunstwerken, Studienobjekten, Büchern und Gelehrtenportraits ausgestattet waren. Diese waren Orte des Sammelns, Studierens und der Generierung von Wissen, vergleichbar mit den Laboratorien in den Naturwissenschaften. Heute findet Forschung in anderen Räumen und mit anderen Mitteln statt. Vernetzte Computer prägen den Arbeitsplatz des Wissenschaftlers. Die Studienobjekte sind digitalisiert und verteilt auf Festplatten, in Datenbanken und im Internet abgelegt und können über unterschiedliche Programme genutzt werden. Für die wissenschaftliche Arbeit droht dadurch ein Ort der Kontinuität verloren zu gehen, an dem die jeweiligen Medienobjekte und Forschungsergebnisse untersucht, arrangiert, miteinander verglichen, bewertet, verknüpft oder kommentiert werden können.

Mit dem Konzept der virtuellen Wissensräume wurde an der Universität Paderborn ein innovativer Ansatz entwickelt, um eine integrierte Nutzung von Medien sowie die Produktion neuer medialer Artefakte im Forschungsprozess oder allgemeiner bei der Wissensarbeit zu ermöglichen (siehe [Ha02] oder [Ke05]). Mediale Objekte werden in unterschiedlichen Konstellationen bearbeitet. Wissensräume enthalten Möglichkeiten zur verteilten Bearbeitung von Dokumenten und zur Diskursstrukturierung ebenso wie Instrumente zum Bewerten und Ordnen von Materialien und zur Koordination gemeinsamer Aktivitäten. Individuen und Gruppen nutzen Wissensräume, die den jeweiligen Wissensstand repräsentieren und als gemeinsames externes Gedächtnis fungieren.

Das Studiolo wird quasi neu erschaffen, und zwar virtuell. Klassische webbasierte Dienste ermöglichen Forschern ebenfalls, Materialien von Servern anderer Forscher herunterladen. Sie sind dann aber bezüglich der weiteren Verarbeitung und Verwaltung lokal auf sich gestellt. Es steht kein gemeinsamer virtueller Arbeitsbereich zur Verfügung, in dem die netzgestützten Materialien verschiedener Forschender und Fachgebiete organisiert, bearbeitet, annotiert und mit anderen diskutiert und ausgetauscht werden können.

Nicht zuletzt die vielen Medienbrüche zwingen dazu, sich den Mechanismen und Konzepten einer verteilten Wissensorganisation zu öffnen. Die Diskussionen um das Web 2.0 signalisieren zudem, dass die Zeit für eine durchgängige Umsetzung gekommen ist und auch solche Funktionen nicht wiederum als isolierte Zusätze betrachtet werden dürfen, indem beispielsweise Wikis oder Blogs als isolierte Einzelanwendungen dem vorhandenen Arsenal hinzugefügt werden.

Virtuelle Wissensräume sind so eine entscheidende Komponente zur Unterstützung der Wissensarbeit. Sie schaffen nicht nur ein gewisses Abbild der realen Welt, sondern ermöglichen sowohl den Erhalt flexibler persönlicher Arbeitsformen als auch innovative Formen der Wissensarbeit. Im Bereich des Lehrens und Lernens wurde das Konzept der Wissensräume u. a. in der Paderborner Lern- und Arbeitsumgebung koALA erfolgreich umgesetzt [Ro07].

3 Wissensarbeit

Im Projekt soll, basierend auf dem Konzept der virtuellen Wissensräume, eine Unterstützungsumgebung für den Bereich Kunst- und Architekturgeschichte aufgebaut werden. Im Zentrum stehen dabei die fachspezifischen Prozesse der Wissensproduktion (siehe Abbildung 1). Im UNESCO Kompetenzzentrum „Materielles und Immaterielles Kulturerbe“, dem Impulsgeber und Pilotanwender der virtuellen Forschungsumgebung, soll damit das Sammeln, Erschließen, Erforschen und Archivieren von erhaltenswertem Kulturgut unterstützt werden. Im Fokus stehen dabei individuelle und kooperative Arbeitsformen sowie die Arrangierbarkeit und Dokumentierbarkeit von Materialien in unterschiedlichen Formaten und aus unterschiedlichen Quellen. Hierzu werden Schnittstellen zu den bereits vorhandenen Bilddatenbanken implementiert und mit den vorhandenen Open-Source-Frameworks für virtuelle Wissensräume Arbeitsbereiche und Funktionen für den Bilddiskurs entwickelt.

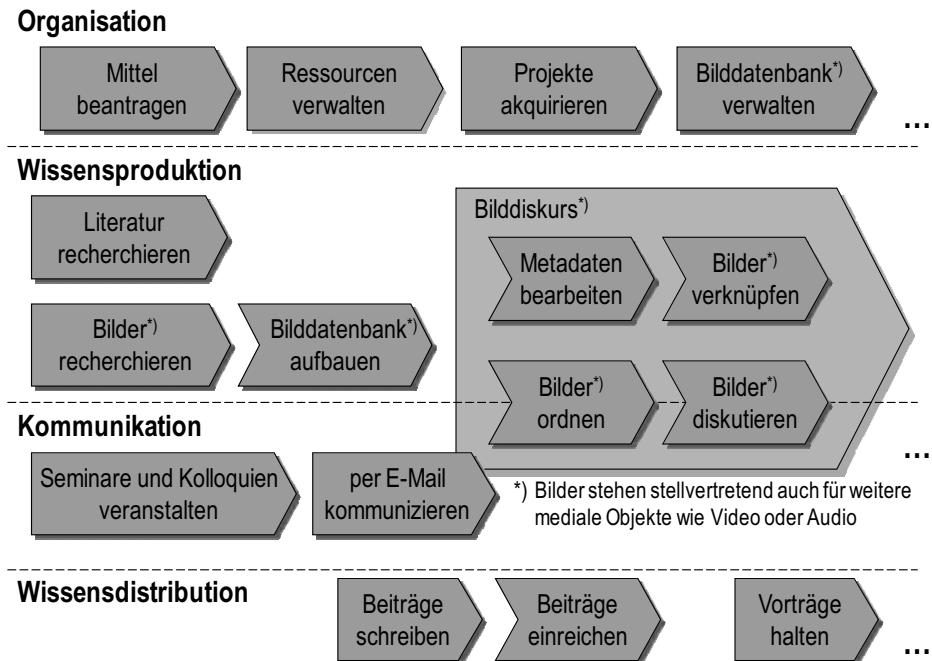


Abbildung 1: Ausschnitt aus der Prozesslandkarte für Forschung in der Kunstgeschichte (Einteilung in Anlehnung an Michael Nentwich [Ne99])

Um den zentralen Forschungsprozess der Wissensproduktion herum gruppieren sich viele Prozesse, die auch in anderen Fachdisziplinen vorhanden und ähnlich gestaltet sind. Im Gegensatz zu der eigentlichen Forschung, zu der Andreas Degkwitz bemerkt, dass sie kein “business“ ist und nicht wie ein „Produktionsbetrieb“ organisiert werden kann [De10], wenden Forscher aber einen Großteil ihrer Arbeitszeit für wenig geliebte Standardaufgaben auf, beispielsweise zur Organisation des Forschungsbetriebs. Diese

Prozesse gilt es zu optimieren und durch geeignete Werkzeuge zu unterstützen, damit möglichst viel Zeit für den kreativen Forschungsprozess bleibt.

4 Standardisierte Dienste zur Forschungsorganisation

Für das IMT als zentralem IT- und Medien-Dienstleister der Universität besteht zusätzlich die Herausforderung, eine virtuelle Forschungsumgebung möglichst wirtschaftlich zu betreiben. Es ist weder personell noch finanziell möglich, für jede Fachdisziplin der Universität eine eigene Forschungsumgebung bereitzustellen. Ziel ist es daher, fachübergreifende Prozesse zu identifizieren und durch geeignete Dienste zu unterstützen und mit fachspezifischen Werkzeugen und Diensten innerhalb der Infrastruktur zu verknüpfen.

Solch ein fachübergreifender Prozess ist das gemeinsame Arbeiten an Dokumenten und die damit einhergehende Dokumentverwaltung. Ein Dienst, der diesen Prozess umsetzt, muss mehr Funktionalität und Komfort bieten, als es bei einem einfachen verteilten Dateisystem der Fall ist.

Als weiterer Dienst wird eine Literaturverwaltung zur Verfügung gestellt. Literatur kann in Katalogen gesucht oder z. B. im BibTex-Format eingeben werden. Die entsprechenden Literaturlisten sind projektbezogen, gehören zu einer Forschungsgruppe oder sind persönlicher Natur.

Die Verwaltung von Forschungsanträgen ist ebenfalls ein Prozess, der sich als zentraler Dienst bereitstellen lässt. Vom Schreiben des Antrags über die Terminierung von Zwischenberichten bis hin zum Abschlussbericht wird der Forscher mit Vorlagen und Workflows unterstützt.

Ein unter anderem durch die Hochschulverwaltung standardisierter Prozess ist die Beschaffung und Verwaltung von Forschungsgeldern und damit einhergehend das Personalmanagement für Mitarbeiter und Hilfskräfte. Der zugehörige Dienst unterstützt mit Vorlagen und Workflows die Beantragung von Mitteln, gibt eine Übersicht über verwendete Ressourcen sowie ablaufende Verträge und stellt die Geldmittel und Arbeitsverträge in Beziehung dar. Die Stundennachweise der Hilfskräfte werden ebenfalls elektronisch erfasst und geben so der Hilfskraft und deren Betreuer eine Sicht auf das Soll und Haben des Arbeitszeitkontos.

Die Publikation von Forschungsergebnissen ist ein weiterer Prozess, der in vielen Disziplinen ähnlich verläuft. Über die vielfältigen Stationen des Prozesses hinweg hilft ein Publikationsdienst dem Forscher, die entsprechenden Informationen und Daten zusammenzuhalten, an die notwendigen Systeme zu übermitteln und dort bei Bedarf zu aktualisieren. Angefangen bei den Forschungsprimärdaten über die ersten Versionen einer Publikation und deren Review bis hin zur Veröffentlichung in einem Bibliothekskatalog stellt der Dienst Funktionen bereit und bindet Schnittstellen an.

5 Infrastrukturen für die Wissenschaft

Bei der Implementierung einer virtuellen Forschungsumgebung muss berücksichtigt werden, dass diese nicht „auf der grünen Wiese“ stattfindet. An den Hochschulen sind bereits vielfältige elektronische Dienste im Einsatz. Einige Dienste können ggf. durch eine bessere Implementierung im Rahmen der virtuellen Forschungsumgebung ersetzt werden. Bei gut eingeführten Diensten sollte jedoch keine Neuimplementierung stattfinden und nach Möglichkeit sollten vorhandene Dienste genutzt werden und die Forschungsumgebung in die vorhandene Infrastruktur integriert werden. An der Universität Paderborn ist eine mehrschichtige Dienste-Infrastruktur [Br10] vorhanden (siehe Abbildung 2).

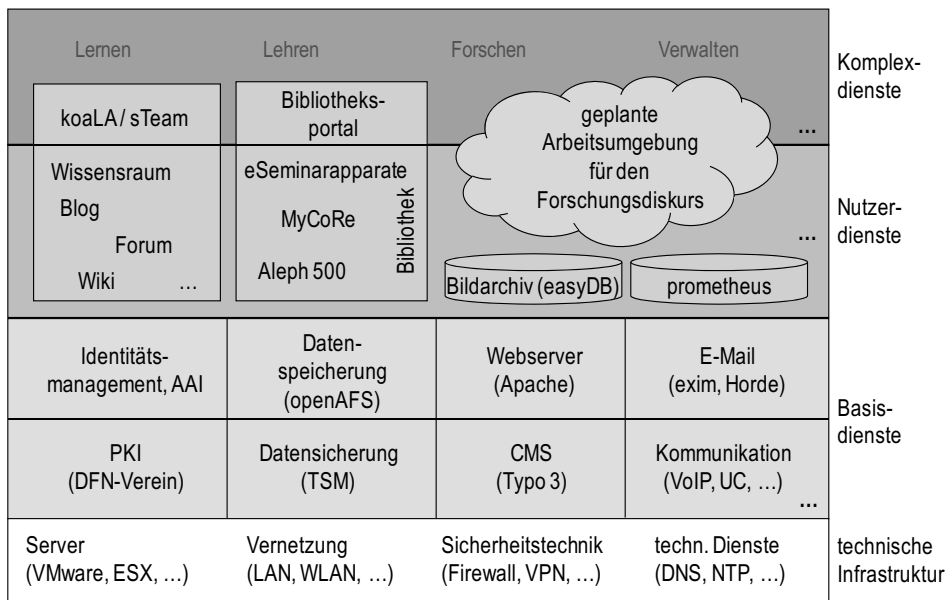


Abbildung 2: Der Aufbau der Paderborner Dienste-Infrastruktur

Die Grundlage bildet die technische Infrastruktur, die beispielsweise die Server, die Vernetzung der einzelnen Rechner sowie erforderliche technische Dienste umfasst. Auf der technischen Infrastruktur aufbauend, übernehmen Basisdienste die kontextunabhängigen Grundfunktionen, die innerhalb der Dienste-Infrastruktur an verschiedenen Stellen immer wieder benötigt werden und die unabhängig von speziellen Aufgaben und Prozessen genutzt werden (bspw. E-Mail). Auf den höheren Schichten sind Nutzerdienste angesiedelt, die bestimmte Nutzungskonstellationen implementieren, während Komplexdienste integrierte Funktionen bereitstellen oder die Integration einzelner Dienste ermöglichen (bspw. über einen Portal-Server). Damit die Dienste der höheren Ebenen angemessen betrieben werden können, müssen sie auf die tieferen Schichten zugreifen können (Authentifizierung, Datensicherung, Netzwerk, ...). Für die Dienste-Infrastruktur ist es essentiell, dass das Fundament stabil und leistungsfähig zur Verfügung steht und über definierte Schnittstellen auch von anderen Anwendungen genutzt werden kann.

6 Möglichkeiten der Umsetzung

Technologisch gibt es zur Realisierung einer virtuellen Forschungsumgebung verschiedene Ansätze. Grundlage für die Umsetzung des virtuellen Arbeitsraumes sollte aus Sicht der Autoren ein Framework sein, das mehreren Ansprüchen genügen muss. Einerseits sollte das Framework viele Standardfunktionen bieten, um unnötige Neuimplementierungen zu vermeiden. Andererseits muss das Framework Möglichkeiten bieten, auf Basis eines komponentenorientierten Ansatzes neue Funktionen bereitzustellen. Die Aufteilung der gesamten Anwendung in Komponenten trägt nicht nur den Methoden der Softwareentwicklung Rechnung, sondern ist auch eine Anforderung des Nutzers, also des Wissenschaftlers. Um den unterschiedlichen Arbeitsweisen in unterschiedlichen Forschungsdisziplinen sowie auch des einzelnen Wissenschaftlers gerecht zu werden, ist die Aufteilung der Anwendung in „Science Apps“ notwendig (vgl. [De10]). Beispielfhaft werden einige im Projekt näher betrachtete Frameworks kurz skizziert.

Ein Open-Source-Framework, das speziell für virtuelle Forschungsumgebungen konzipiert wurde, ist „eSciDoc“ (siehe [Ra09]). Aufbauend auf einem Objekt-Repository (Fedora) werden „Core-Services“ zur Verfügung gestellt, um mit dem Repository auf einem niedrigen Abstraktionsniveau zu interagieren. Darüber befinden sich Services, die speziell für die Anforderungen virtueller Forschungsumgebungen Dienste bereitstellen und dabei auf die einzelnen Core-Services zurückgreifen.¹ Basierend auf eSciDoc wurden mehrere Anwendungen realisiert, so beispielsweise PubMan² (eine ausgereifte Publikationsverwaltung), Faces³ (eine Bilddatenbank für menschliche Emotionen) und ViRR⁴ (eine Kollektion zur Verwaltung und Betrachtung digitalisierter Rechtsschriften). Da eSciDoc ein reines Framework ist, bietet es keinerlei Weboberfläche, die als Basis einer virtuellen Forschungsumgebung eingesetzt werden und als Grundlage einer Weiterentwicklung dienen kann. Betrachtet man den Ansatz von eSciDoc unabhängig von seinem Fokus auf virtuelle Forschungsumgebungen, so zeigt sich im Wesentlichen ein Objekt-Repository mit verschiedenen Services.

Ein weiteres Open-Source-Framework, das dieses Konzept verinnerlicht, ist Alfresco.⁵ Alfresco kann als Enterprise-Content-Management-Anwendung mit kommerziellem Support eingekauft werden sowie als Framework für eigenständige Entwicklungen genutzt werden, da der Quellcode Open Source ist.⁶ Für Alfresco ist eine Weboberfläche verfügbar, die alle Funktionen des Frameworks implementiert und die Anwendung „Alfresco“ darstellt. Diese Weboberfläche kann als Basis für eine Weiterentwicklung genutzt werden.

Auch die kommerzielle Kollaborations-Plattform SharePoint von Microsoft kann als Framework für die Entwicklung individueller Anwendungen dienen. Ein Beispiel für die

¹ vgl. <https://www.esdoc.org/JSPWiki/en/GeneralConcepts>

² <https://www.esdoc.org/JSPWiki/en/PublicationManagement>

³ <https://www.esdoc.org/JSPWiki/en/Faces>

⁴ <https://www.esdoc.org/JSPWiki/en/VirtuellerRaumReichsrecht>

⁵ <http://www.alfresco.com/>

⁶ http://wiki.alfresco.com/wiki/Alfresco_Repository_Architecture

gute Individualisierbarkeit, obwohl SharePoint Closed Source ist, stellt das eLearning-Portal L²P (Lehr- und Lernportal) der RWTH Aachen dar.⁷ Eine eigenständige virtuelle Forschungsumgebung auf Basis von SharePoint ist das Research Information Centre Framework (RIC) [Ba07]⁸. Entwickelt wurde es von Microsoft und der British Library, zunächst mit dem Fokus auf die Biomedizin. SharePoint ist, wie auch Alfresco, eine Webanwendung, in der alle Aspekte des Frameworks bereits implementiert sind. Technologisch bietet SharePoint mit dem Konzept der WebApplications und WebsiteCollections einen Vorteil gegenüber Alfresco. Diese ermöglichen durch eine eigenständige Ressourcen- und Rechteverwaltung das einfache und sichere Hosting autarker Bereiche auf einer zentral bereitgestellten Serverinfrastruktur.

Mit jedem vorgestellten Framework ließe sich eine virtuelle Forschungsumgebung realisieren. Da an der Universität Paderborn und speziell am IMT noch keine Erfahrungen mit einem der vorgestellten Frameworks vorhanden waren, orientierte sich der Entscheidungsprozess nicht ausschließlich an technischen Aspekten, wie der guten Integration in die vorhandene Infrastruktur. Zusätzlich sollte gewährleistet werden, dass der virtuelle Arbeitsplatz von Anfang an eine hohe Akzeptanz bei den Nutzern erlangt.

Die verbreitete Nutzung von Microsoft-Anwendungen, insbesondere Outlook in Verbindung mit Exchange (als E-Mail- und Groupware-Client), Word und PowerPoint haben den Ausschlag dafür gegeben, sich an der Universität Paderborn ausführlicher mit SharePoint zu beschäftigen und basierend auf diesen Erfahrungen SharePoint auch für die Forschungsorganisation einzusetzen, anzupassen und um spezifische Funktionen zu erweitern. Diese Entscheidung wurde auch dadurch beeinflusst, dass am IMT durch die kurz zuvor erfolgreich abgeschlossene Einführung von Microsoft Exchange Erfahrungen im Umgang mit Microsoft-Technologien gesammelt worden waren und von diversen Lehrstühlen und Fakultäten Interesse an einem zentralen SharePoint-Dienst angemeldet wurde.

7 Ein Ansatz zur Forschungsorganisation auf der Basis von SharePoint

SharePoint ist ein Framework, das dem Entwickler nicht bloß die Möglichkeit bietet, Webparts oder Seiten innerhalb der existierenden Oberfläche grafisch zusammenzustellen, sondern ihn mit Hilfe von Visual Studio in die Lage versetzt, komplexe Anwendungen zu realisieren. Das zu Grunde liegende SharePoint-Objekt-Modell ermöglicht es dabei, die Metapher des virtuellen Wissensraums effektiv abzubilden. Arbeitsgruppen bekommen ihren virtuellen Raum, ausgestattet mit individuellen Zugriffsberechtigungen und individueller Funktionalität. Verschiedene Schnittstellen bieten dabei je nach Bedarf unterschiedliche Sichten auf ein und dieselben Daten (Web-Oberfläche, Office-Programme, Explorer).

⁷ siehe <https://www2.elearning.rwth-aachen.de/>

⁸ siehe auch <http://ric.codeplex.com/>

Selbst ohne weitere Entwicklung bietet SharePoint einem Mitarbeiter der Universität Paderborn einen Mehrwert gegenüber der derzeitigen Infrastruktur. Das bestehende Netzwerk-Dateisystem AFS funktioniert zwar zuverlässig, erfüllt aber nicht mehr die Anforderungen eines heutigen Wissenschaftlers. Dokumentsynchronisation muss beispielsweise mit zusätzlichen Werkzeugen durchgeführt werden, wenn man unterwegs (offline) an seinem Notebook arbeiten möchte. SharePoint bietet mit seinen Dokumentbibliotheken hier die Möglichkeit, Dateien einerseits zentral zu speichern und andererseits diese Dokumente lokal für Offlinearbeit verfügbar zu machen, um sie später zu synchronisieren.

Ein weiteres großes Problem beim Arbeiten mit Dokumenten ist die Versionsverwaltung. Neue beziehungsweise geänderte Dokumente werden häufig per E-Mail für einen Review verschickt. Die Verwaltung der unterschiedlichen Fassungen erfolgt dabei in der Regel manuell. Auch hier hilft die zentrale Dokumentbibliothek durch Versionsverwaltung und Check-In/Check-Out-Mechanismen.

Trotz einer tiefen Integration von SharePoint in die Microsoft-Produkte lassen sich die Dokumentbibliotheken auch ohne Office-Produkte von Microsoft nutzen, und durch die webbasierte Umsetzung von SharePoint ist die Dokumentbibliothek auch per Browser erreichbar. Mit der aktuellen Version SharePoint 2010 wurde die Browserkompatibilität verbessert, sodass fast alle Funktionen beispielsweise mit dem Firefox- oder Safari-Browser genutzt werden können. Diese Kompatibilität ist zwingend notwendig, um der heterogenen Systemlandschaft an einer Universität gerecht zu werden.

Über die Standard-Funktionalität hinaus entwickelt das IMT Konzepte, um den Wissenschaftlern Hilfsmittel für die tägliche Arbeit anzubieten. Beim derzeit implementierten Ansatz werden die Wissenschaftler mit drei unterschiedlichen Typen von Arbeitsräumen unterstützt. Auf der Fachgruppenseite, die jedem Lehrstuhl zur Verfügung gestellt wird, existieren neben Verknüpfungen zu den speziellen Forschertools (Bilddatenbank, Publikationsverwaltung etc.) die „Science Apps“. In SharePoint sind dies so genannte Webparts. Für die organisatorischen und generischen Aufgaben gibt es Webparts beispielsweise für die Verwaltung von universitären Accounts der Gäste eines Forschers, die Auflistung von Hilfskräften und deren Vertragsinformationen (Status, Laufzeit, Abrechnungsobjekt) oder die Verwaltung von Austauschprogrammen mit Partneruniversitäten. Für die Wissensdistribution (siehe Abbildung 1) bietet SharePoint unter anderem eine Folienbibliothek an, mit deren Hilfe kooperativ einzelne PowerPoint-Folien erstellt werden können. Für einen konkreten Vortrag wird daraus dann eine individuelle Präsentation zusammengestellt.

Zusätzlich zur Fachgruppenseite stehen Projektseiten als Arbeitsräume zur Verfügung. Benötigt eine Gruppe von Wissenschaftlern einen projektbezogenen Arbeitsraum unabhängig von ihren Fachgruppenseiten, so kann von einem Mitglied der Gruppe selbständig eine Projektseite eingerichtet werden. Je nach Bedarf stehen dann Listen für öffentliche und interne Dokumente, gemeinsame Kalender und Aufgabenlisten bereit. Die Projektseite kann von den beteiligten Fachgruppen aus verlinkt werden.

Den dritten Arbeitsraumtyp bildet die persönliche „MySite“ mit Web 2.0-Funktionen. Hier kann sich der Wissenschaftler selbst darstellen und beispielsweise einen Blog schreiben oder mit seinen Kollegen interagieren. Weiterhin kann hier eine private Dokumentbibliothek gepflegt werden, um beispielsweise Entwürfe für Dokumente zunächst privat zu verfassen, um sie möglicherweise später in die Bibliothek der Fachgruppe zu überführen.

Ein wesentliches Merkmal der SharePoint-Infrastruktur ist der modulare Aufbau. Alle Funktionen und Webparts lassen sich für die Kontexte aktivieren, in denen sie tatsächlich benötigt werden. Dies hilft, die Übersichtlichkeit und damit auch die Akzeptanz durch den Wissenschaftler zu erhöhen. Zugleich wird der Wartungsaufwand für das IMT minimiert, denn wiederkehrende Funktionalität wird zentral entwickelt, bereitgestellt und aktualisiert.

Die Anmeldung an SharePoint erfolgt mit dem Paderborner Uni-Account, der für Studierende, Lehrende und Mitarbeiter im Identitätsmanagement-System verwaltet wird. Für externe Wissenschaftler können über die Gästeverwaltung ebenfalls Accounts in diesem System erstellt werden. Geplant ist hier, diesen Vorgang durch ein föderatives Identitätsmanagement, wie es beispielsweise der DFN-Verein bereitstellt⁹, abzulösen. Dadurch werden Wissenschaftler anderer Einrichtungen in die Lage versetzt, mit ihren eigenen Zugangsdaten kooperativ in SharePoint zu arbeiten. Die Einrichtung eines weiteren Benutzeraccounts an der Universität Paderborn würde dadurch entfallen.

8 Zusammenfassung und Ausblick

Im Projekt wird eine virtuelle Forschungsumgebung realisiert, die als spezifischen Forschungsgegenstand das Arbeiten mit Bildern und anderen multimedialen Dokumenten und darüber hinaus Prozesse um den eigentlichen Forschungsbereich herum unterstützt. Die dazu gewählte Architektur besteht aus einer Kombination unterschiedlicher Frameworks. Für die spezifischen Anforderungen des Bilddiskurses werden Open-Source-Technologien eingesetzt, weil die Notwendigkeit besteht, vorhandene Werkzeuge und Schnittstellen einzubinden, und ein hohes Maß an Flexibilität bei der Implementierung gefordert ist. Für die stärker standardisierbaren Aufgaben im Forschungsprozess wird das kommerzielle Produkt SharePoint eingesetzt, das bereits eine umfangreiche Funktionalität zur Dokumentenverwaltung bereitstellt.

Derzeit werden für die Wissensarbeit mit SharePoint drei unterschiedliche „Arbeitsbereiche“ abgebildet: „MySite“ für die individuelle Arbeit, Bereiche für Lehrstühle zur Organisation des Forschungsbetriebs und der gemeinsamen Verwaltung von Dokumenten sowie Projektbereiche, an denen beispielsweise auch externe Wissenschaftler partizipieren können.

⁹ siehe <https://www.aai.dfn.de/>

Prinzipiell sind die Konzepte übertragbar und auch mit anderen Technologien implementierbar und es ist offen, wie sich die Open-Source-Landschaft in diesem Bereich weiterentwickelt. Im Moment bietet SharePoint durch die Möglichkeiten zur einfachen grafischen Zusammenstellung von Webbereichen jedoch einige Vorteile. Entscheidend für die Akzeptanz an einer Universität sind insbesondere auch das Zusammenspiel mit der vorhandenen heterogenen Systemlandschaft, die Anbindung vorhandener Werkzeuge über Schnittstellen sowie die betriebssystemunabhängige Nutzbarkeit durch die Anwender.

9 Literaturverzeichnis

- [Ba07] Barga, R. S., Andrews, S., Parastatidis, S.: A Virtual Research Environment (VRE) for Bioscience Researchers. In: Proc. of the International Conference on Advanced Engineering Computing and Applications in Sciences (ADVCOMP). IEEE, 2007; S. 31-38.
- [Br10] Brennecke, A., Finke, S., Oevel, G., Roth, A.: Dienste-Infrastrukturen für eLearning – Konzeption, Aufbau und Betrieb. In (Hauenschild, W., Meister, D. M., Schäfer, W. Hrsg.): Hochschulentwicklung innovativ gestalten – Das Projekt Locomotion an der Universität Paderborn. Waxmann, Münster, 2010; S. 41-56.
- [De10] Degkwitz, A.: “Welcome to Science Apps”. Bibliothek: Forschung und Praxis, Preprint (2010).
- [Ha02] Hampel, T.: Virtuelle Wissensräume – ein Ansatz für die kooperative Wissensorganisation. Dissertation, Universität Paderborn 2002 (Online verfügbar: <http://ubdok.uni-paderborn.de/>).
- [Ke05] Keil-Slawik, R.: Dienste-Infrastrukturen als Mittel der Wissensorganisation. In (Kerres, M., Keil-Slawik, R. Hrsg.): Hochschulen im digitalen Zeitalter: Innovationspotenziale und Strukturwandel. education quality forum 2004. Waxmann, Münster, 2005; S. 13-28.
- [Ne99] Nentwich, M.: Cyberscience: Die Zukunft der Wissenschaft im Zeitalter der Informations- und Kommunikationstechnologien; MPIfG Working Paper 99/6, Österreichische Akademie der Wissenschaften; Institut für Technikfolgen-Abschätzung, Wien und Max-Planck-Institut für Gesellschaftsforschung, Köln, Mai 1999.
- [Ra09] Razum, M., Schwichtenberg, F., Wagner, S., Hoppe, M.: eSciDoc Infrastructure – A Fedora-Based e-Research Framework. In (Agosti, M., Borbinha, J., Kapidakis, S., Papa-theodorou, C., Tsakonas G. Hrsg.): Proceedings of the 13th European Conference on Digital Libraries (ECDL 2009), LNCS 5714. Springer, Berlin Heidelberg, 2009; S. 227-238.
- [Ro07] Roth, A., Sprotte, R., Hampel, T., Buese, D.: koaLA – Integrierte Lern- und Arbeitswelten für die Universität 2.0. In (Eibl, C., Magenheimer, J., Schubert, S., Wessner, M. Hrsg.): DeLFI 2007: Die 5. e-Learning Fachtagung Informatik; 17.-20. September 2007 an der Universität Siegen. Gesellschaft für Informatik, Bonn, 2007; S. 221-232.
- [St09] Studiolo communis – Eine ko-aktive Arbeitsumgebung für einen erweiterten Forschungsdiskurs in der Kunst- und Architekturgeschichte. Projektantrag der Universität Paderborn im Rahmen der Ausschreibung „Virtuelle Forschungsumgebungen der DFG, 2009.

Das Projekt „studiolo communis – Aufbau einer ko-aktiven Arbeitsumgebung für einen erweiterten Forschungsdiskurs in der Kunst- und Architekturgeschichte zur Unterstützung des UNESCO Kompetenzzentrums „Materielles und Immaterielles Kulturerbe““ wird im Rahmen des Programms „Wissenschaftliche Literaturversorgungs- und Informationssysteme (LIS)“ von der Deutschen Forschungsgemeinschaft (DFG) gefördert.

Cloud Computing

Zugang zu Föderationen aus Speicher-Clouds mit Hilfe von Shibboleth und WebDAV

Sebastian Rieger¹, Yang Xiang², Harald Richter³

¹Karlsruher Institut für Technologie (KIT),

²Rechenzentrum Garching (RZG),

³Technische Universität Clausthal (TUC)

sebastian.rieger@kit.edu, yang.xiang@rzg.mpg.de, hri@tu-clausthal.de

Abstract: Innerhalb weniger Jahre entstanden sowohl in offenen als auch in geschlossenen Clouds günstige Zugriffsmöglichkeiten auf große Online-Festplatten über das Internet. Sobald jedoch die Benutzer dieser Speicher unterschiedliche Cloud-Dienstleister z.B. für institutsübergreifende Projekte oder für mehrfach verteilte Sicherheitskopien verwenden wollen, treten aufgrund der verschiedenen Benutzerkonten, Zugriffsverfahren und Zugriffsrechte Schwierigkeiten auf. Eine dienstleisterübergreifende, einheitliche Authentifizierung und Autorisierung unter Verwendung einer einzigen Benutzererkennung und eines Passworts wäre für die Benutzer der Cloud-Speicher und für deren Betreiber besser. Der vorliegende Beitrag beschreibt eine Lösung dieses Problems, die aufgrund der Verwendung eines offenen Standards (WebDAV) für den Zugriff auf Online-Filesysteme ohne zusätzliche Middleware auskommt. Die Lösung ist Shibboleth-fähig und damit kompatibel zu einem weit verbreiteten Mechanismus für verteilte Authentifizierung und Autorisierung. Sie beruht auf einer automatischen Benutzer-Lokalisierung im Internet unter Zuhilfenahme von NAPTR-Einträgen im Domain Name System und der Verwendung der E-Mail-Adresse als weltweit eindeutigen Benutzernamen. Die vorgeschlagene Lösung eignet sich für die Realisierung von Föderationen von Speicher-Clouds in denen mehrere Organisationen, z.B. Universitäten und Forschungsinstitute gemeinsam einen einheitlichen Zugriff auf mehrfach verteilte Dateisysteme im Internet bereitstellen wollen, wie sie beispielsweise für die Föderation des Landes Niedersachsen [NAAI] sowie der Max-Planck-Gesellschaft [MAAI] geplant sind.

1 Stand der Technik

In den nachfolgenden Abschnitten dieses Beitrags werden in Kapitel 1 Föderationen von Speicher-Clouds, der Zugang dazu, sowie die Authentifizierung und Autorisierung und das Problem der Lokalisierung von Benutzern beschrieben. Im anschließenden Kapitel 2 wird die Softwarearchitektur unseres Vorschlags, inkl. der Lokalisierung von Benutzern in dynamischen Föderationen dargestellt. In Kapitel 3 schließlich wird ein Fazit gezogen und zukünftige Forschungsarbeiten angegeben.

1.1 Isolierte Speicher-Clouds

Stand der Technik bei Speicher-Clouds sind isolierte Insellösungen. Die Speicher-Cloud lässt sich vom Benutzer über eine Web-Oberfläche oder über eine proprietäre Zugangsschnittstelle wie eine Online-Festplatte mit Verzeichnisstrukturen (häufig auch als sog. Buckets bezeichnet) und Dateien verwenden. Beispiele aus der jüngsten Vergangenheit für solche Speicher-Clouds sind Amazon S3 [S3], Google Storage [GS] und Microsoft Azure Storage [Azure]. Um isolierte Clouds herum sind zusätzlich Dienste und Dienstleister entstanden, die den Zugang zu und die Benutzung von Speicher-Clouds für Endanwender vereinfachen, wie z.B. DropBox [DB], Mozy [MZ] oder Ubuntu One [UO]. Ein einheitlicher de-facto oder de-jure Standard für einen dienstleisterübergreifenden Zugang auf Speicher-Clouds existiert nicht. Es wird aber bei dem Firmenkonsortium SNIA (Storage Networking Industry Association, [SNIA]) unter der Bezeichnung CDMI [CDMI] daran gearbeitet. Wann dieser Industriestandard verfügbar sein wird, ist nicht bekannt.

Neben diesen offenen Speicher-Clouds werden gemäß [Frei10] bei IT-Infrastrukturen im wissenschaftlichen Umfeld zunehmend geschlossene Clouds eingesetzt. Letztere basieren oftmals auf Open Source Implementierungen von Speicher-Clouds (vgl. Eucalyptus Walrus [WALR]) und zeichnen sich dadurch aus, dass den Anwendern innerhalb einer geschlossenen Benutzergruppe ein einheitlicher und vereinfachter Zugriff ermöglicht wird, der instituts- und damit ortsunabhängig ist. Für die Realisierung des Zugriffs erfordert dies i.d.R. eine einheitliche Authentifizierung und Autorisierung (AA) über proprietäre Anwendungen. Ein Beispiel dafür ist die AA-Infrastruktur (AAI) des DFN [DAAI]. Die technische Grundlage für eine AAI bildet häufig der SAML-Standard [SAML] dessen Implementierung in Form von Shibboleth [Mor04] insbesondere in wissenschaftlichen IT-Infrastrukturen weit verbreitet ist.

1.2 Föderationen von Speicher-Clouds

Die natürliche Erweiterung von isolierten Speicher-Clouds sind Verbünde daraus. In [Xia02] wurde eine entsprechende Föderation von Speicher-Clouds beschrieben, die auf REST [REST] beruht. REST basiert wiederum auf HTTP, es verwendet aber u.a. Uniform Resource Identifiers (URIs) anstelle von Uniform Resource Locators (URLs) und neben HTML auch XML in der Response. Der Nachteil dieser Lösung war jedoch, dass Client-seitig zusätzliche Middleware für den Zugriff auf den im Netz verfügbaren Online-Filesystemen erforderlich war.

Das vorliegende Paper beschreibt einen neuen Ansatz, um ohne zusätzliche Middleware auf Föderationen von Speicher-Clouds zugreifen zu können. Dazu wird der WebDAV-Standard [rfc4918] verwendet, der ebenfalls eine Erweiterung von HTTP darstellt, die aber substantiell über REST hinausgeht. Für die vorgeschlagene Lösung wurde ein Shibboleth-fähiger WebDAV-Client entwickelt, der anstelle einer Web-Browser-basierten Benutzerlokalisierung, sowie einer statischen Authentifizierung und Autorisierung (AA) eine sog. dynamische Föderation verwendet, wie sie in [Xia02] beschrieben wurde.

Dadurch können Endanwender auch in zeitlich veränderlichen Föderationen direkt auf unterschiedliche Cloud-Betreiber zugreifen, ohne eine gesonderte Middleware zu verwenden und ohne sich erneut anmelden zu müssen (= Single Sign-On). Die hier vorgeschlagene Lösung erlaubt, das native WebDAV-Modul `mod_dav` [Mdav] des Apache Web-Servers eines Shibboleth-unterstützenden Cloud-Dienstleisters als WebDAV-Server einzusetzen, ohne dass dafür Erweiterungen auf der Server-Seite oder ein Web-Browser für die Verwendung von Shibboleth auf der Client-Seite notwendig wären.

Andere Forschergruppen wie [BeB06] und [Sch08] arbeiten ebenfalls an der Unterstützung von Shibboleth für WebDAV, allerdings auf Seiten des Servers. Eine frühere, Web-Browser-basierte Lösung wurde bereits in [NA07] vorgestellt. Für Grid-Umgebungen wurde eine ähnliche, auf iRODS aufsetzende Lösung in [iRODS] beschrieben.

1.3 Zugang zum Online-Filesystem

Dienstleister für offene Speicher-Clouds haben i.d.R. eine proprietäre Zugangsschnittstelle (vgl. Amazon S3 oder Google Storage REST API), die nur über spezielle Clients genutzt werden kann. Durch letztere werden alle Dateizugriffe auf die bekannten HTTP-Methoden PUT, GET, POST und DELETE abgebildet (vgl. [Xia02]). Einige Cloud Storage Provider bieten darüber hinaus auch einen WebDAV-basierten Zugriff an. Dieser ermöglicht den Benutzern das Lesen, Schreiben und Löschen von Dateien ohne die Verwendung eines speziellen Clients oder einer proprietären API. Dies dient unserer Lösung als Vorbild. Abbildung 1 zeigt, wie der Zugang zweier Beispielanwender zu WebDAV-basierten Online-Filesystemen erfolgt, die in unterschiedlichen Speicher-Clouds angesiedelt sind.

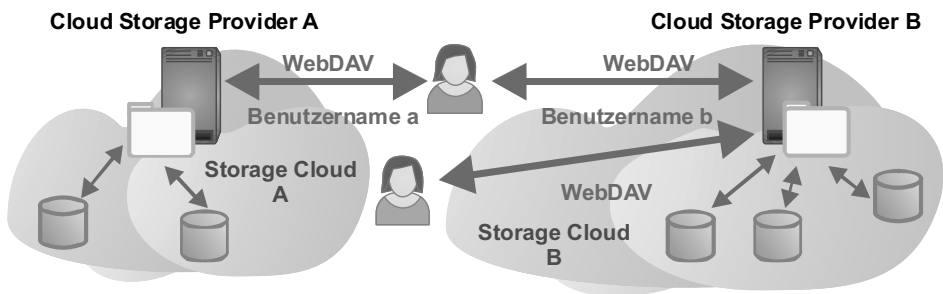


Abbildung 1: Zugang zu WebDAV-basierten Online-Filesystemen in unterschiedlichen Speicher-Clouds.

Die AA erfolgt typischerweise mit Hilfe von Benutzernamen und Passwort. Die Verwendung von Zertifikaten, PINs oder TANs sind im Anwender-Bereich eher unüblich. Für unterschiedliche Cloud Storage Provider müssen Benutzer jedoch i.d.R. verschiedene Benutzernamen und Passwörter verwalten, da jeder Cloud-Betreiber seine Benutzer separat verwaltet und ggf. unterschiedliche Anforderungen an die verwendeten Passwörter stellt. Außerdem ist in Bezug auf die aktuell verfügbaren Anbieter kein Single Sign-On über Cloud-Grenzen hinweg möglich.

Darüber hinaus basiert die Autorisierung von Lese- und Schreibzugriffen bei WebDAV auf den im Kontext des Web-Servers limitierten Zugriffsmöglichkeiten, sowie auf den Rechten, die im darunterliegenden Dateisystem definiert sind. Diese Rechte werden hierbei nicht auf die Benutzernamen, sondern auf eine eindeutige Identifikation (ID) des Benutzers abgebildet. In Windows-Umgebungen bildet beispielsweise der Security Identifier (kurz SID) eine solche ID [SID]. Unix verwendet gemäß des POSIX Standards stattdessen den User Identifier [UID], welcher im Gegensatz zum SID nicht global eindeutig ist. Bei der gleichzeitigen Benutzung mehrerer Speicher-Clouds entsteht daher in Bezug auf die Abbildung der Benutzernamen auf die ID für die Betreiber der Speicher-Cloud ein Problem, das beispielsweise durch die vollständige Verlagerung der Autorisierung in den Web- bzw. WebDAV-Server gelöst werden kann.

1.4 Authentifizierung und Autorisierung in Föderationen

Bei SAML-basierten [SAML] Föderationen benutzen die Diensteanbieter (Service Provider, SP) für die AA sog. Identity Provider (IdP) und lagern die Aufgabe, den Benutzer zu authentifizieren, komplett in den IdP aus. Ein prominentes Beispiel für eine Föderation im wissenschaftlichen Umfeld bildet die AA-Infrastruktur des DFN (= DFN-AAI) [DAAI]. Betrachten wir zur Erläuterung der DFN-AAI einen Mitarbeiter der Max-Planck-Gesellschaft. Dieser wird über den für ihn zuständigen IdP authentifiziert und autorisiert. Der IdP ist an seinem Heimatinstitut angesiedelt, innerhalb dessen der Benutzer einen eindeutigen Benutzernamen und eine wohldefinierte Menge an Zugriffsrechten hat. Aufgrund der Mitgliedschaft der MPG in der DFN-AAI kann der Benutzer jedoch unter Verwendung seines Benutzernamens auch auf Dienstleistungen zuzugreifen, die von Instituten außerhalb der MPG angeboten werden, sofern diese Institute ebenfalls an der DFN-AAI teilnehmen. Da in einer Föderation nur der für den Benutzer zuständige IdP die AA durchführt, wird Single Sign-On auch für den Spezialfall einer Föderation über Cloud-Grenzen hinweg möglich. Cloud Service Provider (CSP) einer Speicher-Cloud, verwenden den jeweils zuständigen IdP einer Föderation für die AA. Diese Methode ist in Abbildung 2 dargestellt.

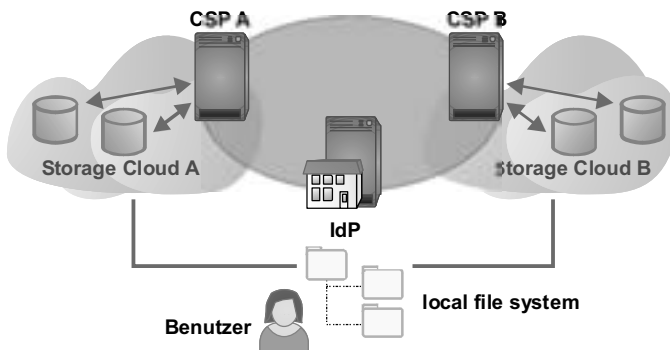


Abbildung 2: Single Sign-On in Föderationen aus Speicher-Clouds mit Hilfe eines Identity Providers (IdP).

Aufgrund des Verbundes mehrerer Partner und aufgrund besserer Skalierbarkeit existieren in einer Föderation mehrere IdPs. Diese müssen innerhalb der Föderation und bei den einzelnen CSPs in der Speicher-Cloud verwaltet werden. Um diese Verwaltung zu vereinfachen, kann in einer Cloud ein Identity-Dienst, dessen Funktion als „Identity as a Service“ (IDaaS) bezeichnet wird, verwendet werden. In dem Szenario aus CSP, IdPs und zentralem Identity-Dienst bildet der Identity-Dienst mit den IdPs eine sternförmige Topologie und wirkt für die CSPs als Proxy der IdPs. Der Vorteil, der sich für die CSPs aus der sternförmigen Topologie ergibt, ist die Vereinfachung der AA, da nur zum zentralen Identity-Dienst ein Vertrauensverhältnis bestehen muss. Insgesamt entsteht eine mehrfache Indirektion des Vertrauens gemäß des Transitivitätsgesetzes, beginnend mit dem Vertrauensverhältnis vom CSP zum zentralen Identity-Dienst, von dort zu den einzelnen IdPs und schließlich zu den Benutzern. Diese Indirektion ist in Abbildung 3 dargestellt und wurde in [Xia02] beschrieben.

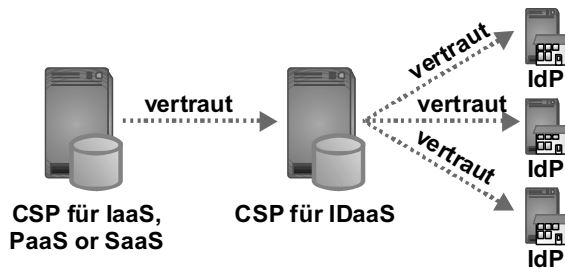


Abbildung 3: Mehrfache Indirektion des Vertrauens für AA in einer Föderation gemäß [Xia02].

1.5 Lokalisierung des Benutzers im Internet

Im Folgenden wird davon ausgegangen, dass der SAML-basierte Shibboleth Standard für die AA verwendet wird. Dieser bildet auch die Basis für unsere in Kapitel 2 vorgeschlagene Lösung für den föderativen Zugriff auf Speicher-Clouds. Shibboleth verwendet für die Lokalisierung des Benutzers, d.h. für die Ermittlung von dessen Heimatorganisation einen sog. Discovery Service (DS). Der DS stellt eine Web-Seite mit fester und a priori bekannter URL bereit, auf die der Benutzer zugreift. Während des Zugriffs auf den SP wird der Benutzer auf diese Web-Seite umgeleitet. Auf der Anmelde-Web-Seite wird dem Benutzer eine Liste von Organisationen bzw. IdPs präsentiert, die an der Föderation teilnehmen. Der Benutzer wählt seine Heimatorganisation und damit auch den für ihn zuständigen IdP. Dann wird der Web-Browser des Benutzers über eine HTTP Redirection an denjenigen IdP weitergeleitet, bei dem er seinen Benutzernamen und Passwort eingeben muss, um die Authentifizierung durchzuführen. Sind mehrere Föderationen z.B. unter Verwendung von eduGAIN [EduG] zu einer sog. Konföderation zusammengefasst, ähnlich wie dies auch bei der bereits skizzierten Föderation von Speicher-Clouds der Fall ist, so werden die Discovery Services der Konföderation sowie der darin enthaltenen Föderationen kaskadiert, d.h., der Benutzer meldet sich im Sinne einer Anmeldehierarchie an.

Der Benutzer wählt hierbei zunächst im DS der Konföderation seine Föderation aus. Anschließend wird er an den DS dieser Föderation umgeleitet, bei dem er seine Heimatorganisation resp. seinen IdP selektiert. Diese Auswahl ist für einen WebDAV-Zugriff ohne Web-Browser nicht praktikabel. Der Benutzer wünscht sich hier, wie in Abbildung 2 dargestellt, eine direkte Anbindung in Form eines virtuellen Dateisystems, ohne beim Zugriff auf Verzeichnisse und Dateien zusätzlich einen Web-Browser öffnen zu müssen, bzw. zwischen Datei-Explorer und Web-Browser zu wechseln. Eine erneute separate Anmeldung für jeden Storage Provider, z.B. bei einer Aggregation mehrerer Dateisysteme über mehrere CSPs, wäre ebenfalls benutzerunfreundlich.

Zusätzlich versagt diese Methode, wenn sich die Zahl der IdPs in einer Speicher-Cloud schnell ändert, oder wenn die Zahl der Speicher-Clouds in einer Föderation zeitlich variiert. Für diesen komplizierteren Fall kann die AA über eine in [Xia01] beschriebene dynamischen Discovery-Prozedur erfolgen: Zuerst wählt der Benutzer, wie bereits beschrieben, den CSP an. Danach gibt er jedoch anstelle der Auswahl seiner Heimatorganisation als Benutzernamen direkt seine E-Mail-Adresse ein, und ein erweiterter Discovery Service sendet basierend auf der Domain der E-Mail Adresse eine DNS-Anfrage an den zuständigen DNS-Server. Das Ergebnis der Anfrage ist ein DNS NAPTR-Eintrag, in dem zuvor der für den Benutzer verantwortliche IdP verbucht wurde. Der NAPTR-Eintrag wird dann vom Cloud-Dienstleister ausgewertet und mit dessen Information der richtige IdP konsultiert, der die AA durchführt ([Xia02]). Durch die dynamische Discovery-Prozedur können Benutzer auch über die Grenzen einer Föderation hinweg lokalisiert werden. Dies wird in [Xia02] auch als dynamische Föderation bezeichnet.

Der geschilderte Vorgang ist in Abbildung 4 gezeigt. Der Benutzer erhält nach einmaliger erfolgreicher Anmeldung an seinem IdP über alle SPs innerhalb der dynamischen Föderation ein Single Sign-On.

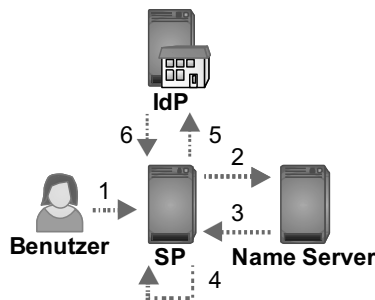


Abbildung 4: AA mit Hilfe eines DNS NAPTR-Eintrags in zeitlich veränderlichen Föderationen.

Das Konzept eignet sich auch für Benutzer, die mehrere Identitäten bei unterschiedlichen IdPs haben. Diese können durch die Angabe Ihrer E-Mail-Adresse selbst entscheiden, welche Identität sie bei der Anmeldung am SP verwenden wollen.

2 Zugang zu Speicher-Clouds mit Shibboleth und WebDAV

In diesem Kapitel wird die Architektur eines Shibboleth-fähigen WebDAV-Clients für den Zugang zu Föderationen von Speicher-Clouds beschrieben.

Der Client basiert auf einer Erweiterung des WebDAV-Clients Sardine [Sard], der als Open Source zur Verfügung steht, und ist als Swing-basierte Java-Anwendung realisiert. Der implementierte Prototyp erlaubt den Zugriff auf Dateiverzeichnisse, deren Inhalten, sowie das Kopieren einzelner Dateien zwischen WebDAV-Server und lokalem Rechner mittels „Drag & Drop“. Sardine benutzt den Apache HTTP Client [HC] für den Zugriff auf den Web-Server. Um eine AA per Shibboleth zu ermöglichen, wurde eine zusätzliche Methode für die Verbindung zum WebDAV-Server realisiert. Diese Methode verarbeitet die für Shibboleth erforderlichen Weiterleitungen (HTTP Redirects) bei der AA, sowie die Auswahl des IdPs unter Verwendung der in Kapitel 1.5 "Lokalisierung des Benutzers im Internet" erläuterten dynamischen Discovery-Prozedur. Dabei entfällt die manuelle Auswahl des IdPs und wird durch die automatisierte Ermittlung des zuständigen IdPs anhand der E-Mail-Adresse ersetzt. Die Methode extrahiert die SAML-Antwort (inkl. Relay State), die der IdP nach der erfolgreichen Authentifizierung des Benutzers erzeugt, und sendet diese (per SAML HTTP POST Profile) an den CSP zurück. Ferner wurde Sardine von uns um eine Sitzungsverwaltung erweitert, die die vom SP und vom IdP erstellten HTTP-Sitzungs-Cookies [rfc2965] während der Verwendung des Clients analog zur Sitzungsverwaltung eines Web-Browsers speichern.

Dadurch wird ein Single Sign-On sowohl über unterschiedliche CSPs als auch über die Grenzen einzelner Clouds hinweg realisiert. Die Autorisierung der Benutzer erfolgt anhand von Attributen, die der IdP an den jeweiligen SP übermittelt. Diese Attribute können in unterschiedlichen im Apache Web-Server konfigurierten sog. Locations, die die WebDAV Verzeichnisstruktur abbilden, oder im SP selbst für die Autorisierung verwendet werden.

Die skizzierte Lösung kann mit gängigen Web-Servern wie z.B. Apache oder Microsoft IIS für den WebDAV-basierten Zugriff auf entfernte Dateien und Verzeichnisse verwendet werden. Innerhalb des für den Test unseres Prototyps verwendeten Apache Web-Servers sind hierfür die Module `mod_dav` und `mod_davfs` vorhanden. Für die Authentifizierung benutzt `mod_dav` wiederum das Apache Modul `mod_auth` und dessen Authentifizierungserweiterungen. Aus der Sicht des Apache Servers bildet der Shibboleth SP eine solche Authentifizierungserweiterung (`mod_shib`). Verwendet man diese für eine Location des Apache Web-Servers, die mittels `mod_dav` für WebDAV eingerichtet wurde, so wird der Benutzer beim ersten Zugriff mit einem Web-Browser auf diese Location automatisch an den Discovery Service und von dort nach entsprechender Auswahl an den IdP seiner Heimatorganisation weitergeleitet. Nach erfolgreicher AA erfolgt eine Umleitung zurück an den SP. Dieser erlaubt, sofern die Zugriffsrechte vorhanden sind, den Zugriff auf die in der Location vorhandenen Verzeichnisse und Dateien. Veränderungen an den Dateien und Verzeichnissen sind mit gängigen Web-Browsern allerdings nicht möglich. Hierfür wird unser WebDAV-Client benutzt.

Um ein Funktionieren auch bei zeitlich veränderlichen Föderationen über Speicher-Clouds zu ermöglichen, verfügt jeder IdP und SP über einen sog. Trust Estimation Service (TES) [Xia01]. Der TES wurde von uns als Erweiterung in Shibboleth integriert. Wie Abbildung 5 zeigt, fungiert der TES für Shibboleth wie ein externer Discovery Service. Der TES ist als standalone Tomcat-Anwendung implementiert und wurde mit den Standard-Shibboleth Komponenten verbunden.

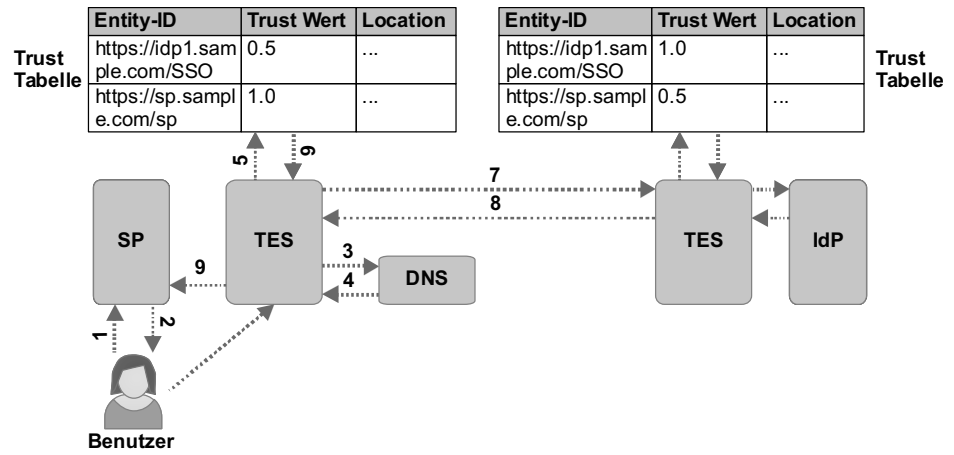


Abbildung 5: Integration eines sog. Trust Estimation Service in die Shibboleth.
 Die Benutzeroberfläche der erreichten Lösung ist in Abbildung 6 dargestellt.

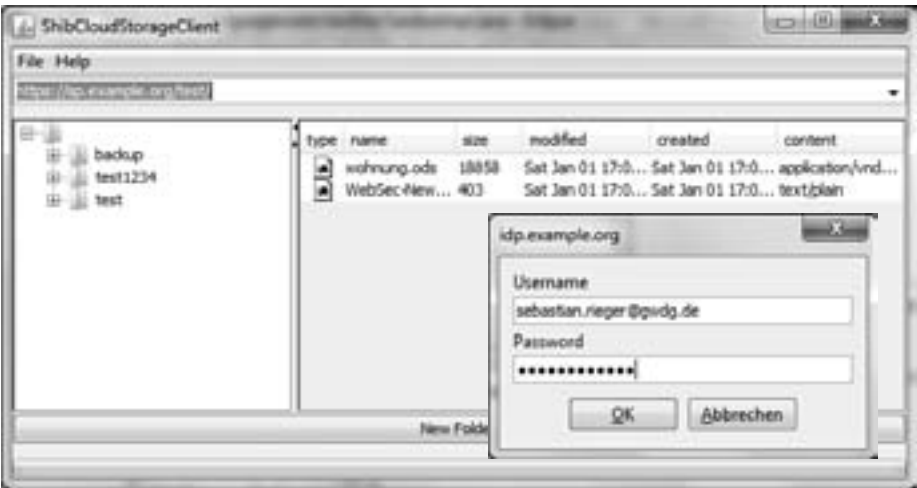


Abbildung 6: Benutzeroberfläche für den Zugang zu Föderationen von Speicher-Clouds.

3 Fazit und Ausblick

Die vorgestellte Lösung ermöglicht eine einheitliche Authentifizierung für Speicher-Clouds. Der implementierte Shibboleth-fähige WebDAV-Client erlaubt dabei durch die Verwendung einer dynamischen Discovery-Prozedur ein Single Sign-On über mehrere Cloud Storage Provider, d.h. mehrere Speicher-Clouds hinweg. Das Lesen und Schreiben auf den Online-Festplatten erfolgt durch die Verwendung von WebDAV ähnlich wie in einem virtuellen Dateisystem. Dies ermöglicht einen einheitlichen Zugriff auf verteilte Dateisysteme für die Realisierung von standortübergreifenden Speicher-Clouds, wie sie beispielsweise innerhalb der MPG-AAI [MPAAI] und Nds-AAI [NAAI] geplant sind. In Zukunft wird evaluiert, wie die Abbildung von Benutzernamen auf eindeutige IDs für die Zugriffsrechte des Dateisystems ohne die explizite Verlagerung der Autorisierung in den Web-Server oder in den Shibboleth-fähigen Service Provider realisiert werden kann. Darüber hinaus soll geprüft werden, inwieweit sich objektbasierte Online-Filesysteme (z.B. auf Basis von NoSQL-Datenbanken) für die Integration in ein virtuelles Dateisystem bei Föderationen von Speicher-Clouds eignen, die für ihre Datenreplikation bereits global eindeutige Identifikation verwenden (vgl. das virtuelle Dateisystem GridFS für MongoDB [Monfs]).

Literaturverzeichnis

- [ACL] Access Control List: http://de.wikipedia.org/wiki/Access_Control_List (6. Januar 2011).
- [Azure] Windows Azure Storage: <http://www.microsoft.com/windowsazure/windowsazure/> (6. Januar 2011).
- [BeB06] Bellembois, T.; Bourges, R.: The open-source ESUP-Portail WebDAV storage solution, <http://www.esup-portail.org/download/attachments/43515911/ESUP-Web-DAV.pdf> (6. Januar 2011).
- [CDMI] Cloud data management interface. SNIA Web Site, April 2010: <http://cdmi.sniacloud.com/> (6. Januar 2011).
- [DAAI] DFN-AAI – Authentifikations- und Autorisierungs-Infrastruktur: <https://www.aai.dfn.de> (6. Januar 2011).
- [DB] Dropbox: <http://www.dropbox.com/> und <http://de.wikipedia.org/wiki/Dropbox> (6. Januar 2011).
- [EduG] eduGAIN: <http://www.edugain.org/> (6. Januar 2011).
- [Frei10] Freitag, S: Erweiterung einer D-Grid-Ressource um eine Compute-Cloud-Schnittstelle. In (Müller, P.; Neumair, B.; Dreo Rodosek, G., Hrsg.): Proc. 3. DFN-Forum Kommunikationstechnologien, Konstanz 2010. Gesellschaft für Informatik, Bonn, 2010; S. 13-22.
- [GS] Google Storage: <http://code.google.com/intl/de-DE/apis/storage/> (6. Januar 2011).
- [HC] Apache HttpComponents: <http://hc.apache.org/> (6. Januar 2011).
- [iRODS] Zhang, S.; Coddington, P.; Wendelborn, A.: Davis: A generic interface for iRODS and SRB, 10th IEEE/ACM International Conference on Grid Computing, Banff, 2009.
- [Mdav] Apache Module mod_dav: http://httpd.apache.org/docs/2.0/mod/mod_dav.html (6. Januar 2011).
- [Monfs] MongoDB GridFS Specification: <http://www.mongodb.org/display/DOCS/GridFS+Specification> (6. Januar 2011).
- [Mor04] Morgan, R. L.; Cantor, S.; Hoehn, W.; Klingenstein, K.: Federated Security: The Shibboleth Approach, EDUCAUSE Quarterly, Vol. 27, 2004, S. 12-17.
- [MAAI] Max-Planck-Gesellschaft - MPG-AAI: <https://aai.mpg.de> (6. Januar 2011).

- [MZ] Mozy: <http://mozy.com> und <http://en.wikipedia.org/wiki/Mozy> (6. Januar 2011).
- [NA07] Ngo, L.; Apon, A.: Using Shibboleth for Authorization and Authentication to the Subversion Version Control Repository System, Fourth International Conference on Information Technology - ITNG '07, Las Vegas, 2007.
- [NAAI] Nds-AAI: Authentifizierungs- und Autorisierungs-Infrastruktur für Niedersachsen: <http://www.daasi.de/projects/ndsaaai.html> (6. Januar 2011).
- [REST] Roy T. Fielding. Architectural Styles and the Design of Network-based Software Architectures. PhD thesis, University of California, Irvine, 2000 und http://de.wikipedia.org/wiki/Representational_State_Transfer (6. Januar 2011).
- [rfc2965] Kristol, D.; Montulli, L.: HTTP State Management Mechanism, <ftp://ftp.rfc-editor.org/in-notes/rfc2965.txt> (6. Januar 2011).
- [rfc4918] Dusseault, L. et al.: HTTP Extensions for Web Distributed Authoring and Versioning (WebDAV), <ftp://ftp.rfc-editor.org/in-notes/rfc4918.txt> (6. Januar 2011).
- [S3] Amazon Simple Storage Service (Amazon S3), <http://aws.amazon.com/de/s3/>, abgerufen am: 6.1.2011.
- [SAML] OASIS: Security Services (SAML) TC: www.oasis-open.org/committees/security/ (6. Januar 2011).
- [Sard] sardine - an easy to use webdav client for java: <http://code.google.com/p/sardine/> (6. Januar 2011).
- [Sch08] DataFinder: A Python Application for Scientific Data Management, EuroPython 2008: The European Python Conference, Vilnius, 2008.
- [SID] Security Identifier: http://de.wikipedia.org/wiki/Security_Identifier (6. Januar 2011).
- [SNIA] Storage Networking Industry Association: <http://www.snia.org> und http://de.wikipedia.org/wiki/Storage_Networking_Industry_Association (6. Januar 2011).
- [UID] Benutzerkennung: <http://de.wikipedia.org/wiki/Benutzerkennung> (6. Januar 2011).
- [UO] Ubuntu one: <https://one.ubuntu.com/> (6. Januar 2011).
- [WALR] Interacting with Walrus (2.0) - Storage Service: http://open.eucalyptus.com/wiki/EucalyptusWalrusInteracting_v2.0 (6. Januar 2011).
- [Xia01] Xiang, Y.; Kennedy, J.A.; Richter, H.; Egger, M.: Network and Trust Model for Dynamic Federation, The Fourth International Conference on Advanced Engineering Computing and Applications in Sciences, Florence, 2010.
- [Xia02] Xiang, Y.; Rieger, S.; Richter, H.: Introducing a Dynamic Federation Model for RESTful Cloud Storage, The First International Conference on Cloud Computing, GRIDs, and Virtualization, Lisbon, 2010.

Polybius: Secure Web Single-Sign-On for Legacy Applications

Pascal Gienger Marcel Waldvogel

University of Konstanz

Konstanz, Germany

pascal.gienger@uni-konstanz.de, marcel.waldvogel@uni-konstanz.de

Abstract: Web-based interfaces to applications in all domains of university life are surging. Given the diverse demands in and the histories of universities, combined with the rapid IT industry developments, all attempts at a sole all-encompassing platform for single-sign-on (SSO) will remain futile. In this paper, we present an architecture for a meta-SSO, which is able to seamlessly integrate with a wide variety of existing local sign-in and SSO mechanisms. It is therefore an excellent candidate for a university-wide all-purpose SSO system. Among the highlights are: No passwords are ever stored on disk, neither in the browser nor in the gateway; its basics have been implemented in a simple, yet versatile Apache module; and it can help reducing the impact of security problems anywhere in the system. It could even form the basis for secure inter-university collaborations and mutual outsourcing.

1 Introduction

The processes in teaching, learning, research, and administration at a university are manifold. As a result, the applications written for their support date back to various decades, using dozens of different interfaces, programming environments, and authentication/authorization concepts. Despite efforts into identity management, the integration in practice is limited to a subset of applications, as it falls short of user demands and security requirements. Outside a small group of “compatible” applications, there remains a chaos of usernames, passwords, and conventions, combined with many manual and error-prone processes. Even those few applications are sometimes combined by hard-to-use SSO systems, turning ordinary users away in despair.

In this paper, we present a generic framework, *Polybius*,¹ which allows the creation of an SSO mechanism which provides single-log-out and can sit on top of existing SSO or local sign-in systems. It avoids the common pitfalls of storing passwords in cleartext anywhere. We also have implemented the system in a light-weight Apache module, providing an interface to three previously unconnected systems, none of which require any modifications to become part of the SSO infrastructure.

¹Polybius was a Greek historian who also reported on the secure way the Romans used to distribute their ephemeral passwords http://en.wikipedia.org/wiki/Password#History_of_passwords (visited 2011-01-06)

The processes at a university are not only manifold, as already explained, but also tightly interwoven. However, the various applications stemming from different decades, technologies, and manufacturers, rarely provide this integration directly.

For example, in the context of a research project, acquisition, funding planning, people hiring, project planning, e-collaboration, and travel require several distinct applications, requiring different tools and multiple entry of the same information. Also, lectures require time and room scheduling, assigning assistants, group assignments, preparation and publishing of texts, notes and slides, as well as grading and entering grades, which again involve a multitude of systems.

Therefore, important goals of application management at universities include:

1. to integrate these applications from a user's point of view, e.g., portals, common authentication infrastructure, ...; and
2. to allow applications to interoperate, i.e., have data generated by one application be used to .

The aim of both types of integration is to improve the efficiency of processes while at the same time allowing modularity to prevent single-vendor lock-ins and therefore provide competitive multi-vendor integration or migrations.

Conventional designs have kept the SSO mechanism out of the actual user $\leftrightarrow$ system data flow as much as possible. However, the time is ripe to rethink traditional design rationale: The secure design of Polybius, advances in computer performance, coupled with ease of redundancy and graceful degradation make Polybius less of a single point of failure than conventional SSOs in many cases. We discuss this in detail in Section 3 and show applications with and without password synchronization in Sections 4 and 5, respectively.

2 Related work

Many different approaches to the SSO problem have been developed, most notably Kerberos [NYHR05], Shibboleth [Sco05], and CAS [Jas].

Kerberos is very widespread, frequently deployed as part of Microsoft Active Directory. Applications do check issued Kerberos tickets which are requested in authentication phase from the Key Distribution Center, KDC ('Domain Controller' in Active Directory parlance). The concept has been slightly modified by Microsoft to be able to send Kerberos tickets through by HTTP to allow their Microsoft Internet Explorer to use the same authentication structure for web applications [JZB06].

In an effort to offer an SSO infrastructure to a broader audience and heterogeneous environments, other schemes did evolve. Shibboleth is widely used in the European academic community whereas CAS is often found in North America. Shibboleth is well known for its cross-domain authentication, which we will not cover here.

All these methods share disadvantages when it comes to legacy web applications²: They will not work out of the box. Instead, you will have to rewrite each and every application to make use of these structures. Applications have to be “kerberized” to use Kerberos, “shibbolethised” to use Shibboleth, and so on.

First, any such modification requires access to the application source code to modify the login process. Second, it may even require modification of the application logic, as different information is available as part of the login process. Third, it may not even work at all, if the backend insists on password-based authentication, as is common when accessing legacy services such as terminal emulations or non-Web protocols including many IMAP or LDAP servers.

3 Polybius design

To avoid these problems, our approach is to keep the passwords as part of the process. Storing or transmitting passwords is generally frowned upon, as they increase the vulnerability of the system: An attacker may gain access to the password store or to the user’s browser. To avoid attacks on data at rest or data in flight, Polybius employs a form of secret sharing between the browser and the SSO gateway: Neither of them has enough information to get at the password, but on every request, the browser delivers the information which allows the gateway to recover the password for this request; which the gateway forgets immediately after using it for the backend service.

To accomplish this, user credentials (e.g. user name(s) and password) are stored in a session database encrypted using AES in CBC mode, using a 256 bit encryption key and a 128 bit initialization vector (IV). These two values together with a session ID (to reference the appropriate record in the database) form the unique session “key” needed by the SSO server to reconstruct the password when needed. This SSO session key is stored as an HTTP cookie in the user’s browser (cf. Figure 1).

This method leaves the session database useless for an attacker – without the keys he is not able to extract the credentials. When capturing a cookie from a browser the attacker will be able to take over the session (as it is the problem with every cookie based session mechanism) but even if he gets the session database he would only be able to decrypt the session record tied to his session key. Compared to other methods (storing passwords in cleartext during a session) this reduces the abilities of an attacker to gain any secrets.

The latter can be minimized in impact by expanding the scheme to random temporary session passwords, created especially for this session only (see below).

Polybius runs as an Apache module and makes use of the reverse proxy features. Every web application running under Polybius SSO is accessed using an application-specific prefix specified in the proxy configuration. The chosen prefix makes Apache forward

²We define a **legacy web application** to be an application which does not use some shared authentication mechanism consisting of exchanging cryptograms like kerberos tickets or shibboleth hashes to be verified against an authentication server and which cannot be modified to use them. Legacy web applications rely on the input of a username and a related password via an HTTP POST request or a WWW-Authenticate-Header.

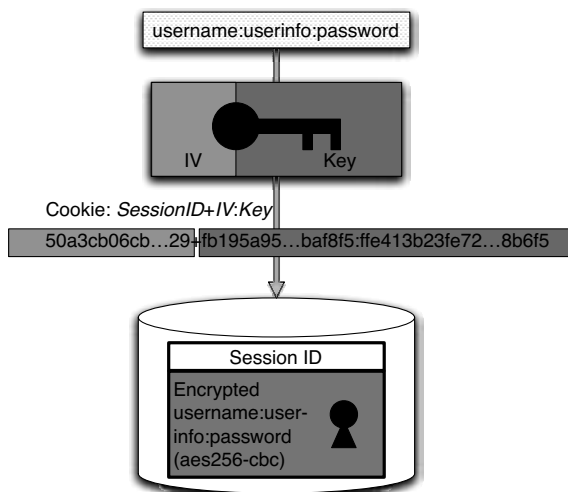


Figure 1: User credentials are stored encrypted in the session database.

(reverse proxy) the request to the specified web application. Apache's proxy modules are capable of rewriting HTTP path responses, cookie domains, and cookie paths.

Applications may still be accessed without a Single-Sign-On. Every application may remain a standalone installation. Polybius also can be used to allow helper applications to automate processes involving one or more backend systems. Single logout is possible by simply erasing the database record belonging to this user.

4 Polybius with synchronized passwords

In this case, the organization has already integrated its applications to access a single user store or multiple stores with synchronized passwords. This can be obtained by identity or user management software.³

4.1 Basic HTTP authentication

The interaction with legacy web applications using HTTP basic authentication⁴ is shown in Figure 2.

³The actual identities need not to be the same, but then Polybius needs access to a service which can link those identities.

⁴WWW-Authenticate: header in the server's reply, followed by the client repeating the request including Authorization: Basic XXXX in the request, where XXXX is the base64-encoded form of username:password.

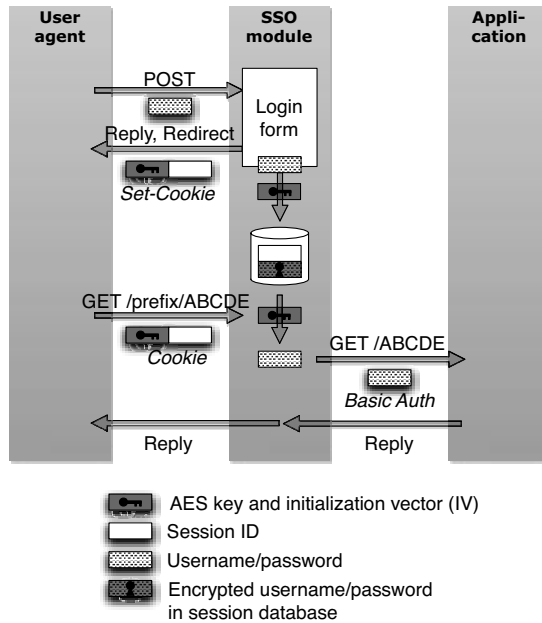


Figure 2: Cookie flow diagram for the basic http auth scenario

- The user enters his username and his password on a HTML Form which is part of Polybius' infrastructure.
- The password has been checked against the authentication database. If it is wrong, the servers aborts the request with an appropriate error.
- Random 256 bit AES key, 128 bit IV, and 128 bit Session-ID are generated.
- The user name(s) and the password are encrypted with this (key, IV) tuple using CBC mode of AES.
- This encrypted data is stored in the session database under the Session ID as the database key.
- A **Set-Cookie**: header is sent in the reply to the user's browser consisting of the session ID, the AES key, and the IV.

The user is now *logged on*. The legacy application is defined as an apache proxy target under a specific prefix. So whenever a GET/POST request arrives regarding an URL under this prefix the following happens:

- The user browser sends the SSO cookie along with his request. The Polybius Apache module extracts this cookie and deletes it from the request header. The application behind the Apache proxy will not see it.

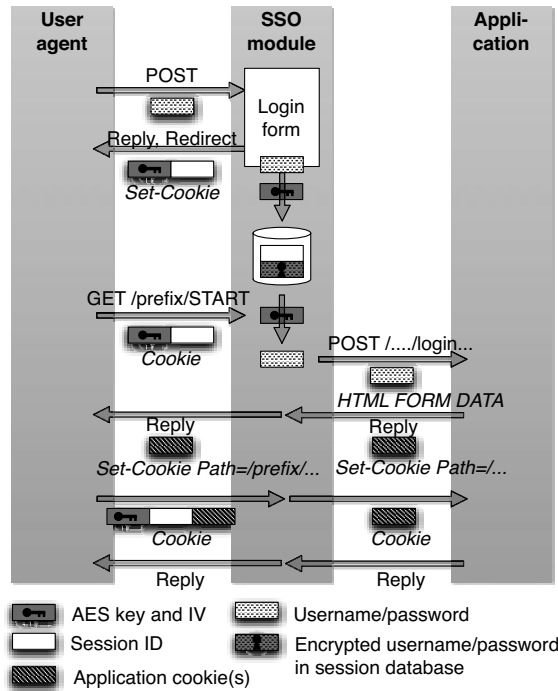


Figure 3: Cookie flow diagram for the session cookie scenario

- The SSO module retrieves the session data stored in the database under the session ID given in the browser cookie. It then decrypts it using the AES key and IV also included in the cookie.
- The SSO module now constructs an `Authentication: Basic XXXX` header line in the request using the username and the password decrypted from the session database.
- The Apache proxy module now takes this HTTP request and forwards it to the web application.
- The response is passed through to the user's browser.

4.2 Cookie-based authentication

The interaction with legacy web applications using a cookie-based session management is also possible using a “Polybius Login Helper” for the application. The communication diagram is shown as Figure 3.

- After the login phase (same as above), the user requests a special session start URL

which in turn starts a login helper for that application.

- This login helper POSTs the appropriate FORM data to the application to “log in” as if it were the real user agent.
- The cookie given back by the application is passed slightly modified to the user’s browser: The cookie path specification is changed so that this cookie is only sent when requesting the application (distinguished by the Apache proxy prefix).
- The SSO module ensures that neither SSO key nor SSO session ID are passed to the application. Also, no secrets from the backend cookie will be leaked to browser.

4.3 Example Apache configuration

To include a service authenticated using HTTP Basic, such as an Subversion (SVN) version-control repository, the following Apache configuration fragment could be used:

```
LoadModule polybius_module /usr/lib64/apache/modules/mod_polybius.so
SetOutputFilter POLYBIUS
SSLProxyEngine On

ProxyPass /svn/ https://svn.uni-konstanz.de/
ProxyPassReverse /svn/ https://svn.uni-konstanz.de/

<Location /svn>
    PolybiusAuthType basic
    PolybiusUidAttribute cfn
</Location>
```

5 Polybius with temporary passwords

Sometimes passwords cannot be synchronized between all applications. This can be due to conflicting character set restrictions or due to the fact that the application does not include a mechanism compatible with the identity management system.

This is no problem for Polybius, as the schema above can be extended to make use of temporary session passwords instead of “real” ones. Some changes in the authentication infrastructure is needed however, as the password (and the respective hashes) is a unique single-value property in many systems. The concept is to slightly modify the existing infrastructure to add multiple temporary passwords⁵ against the applications may authenticate instead of the “real” one. The idea is sketched in Figure 4.

⁵The usage of a single-value would result in a situation where multiple concurrent logins for the same user are prohibited - which can be a desired feature.

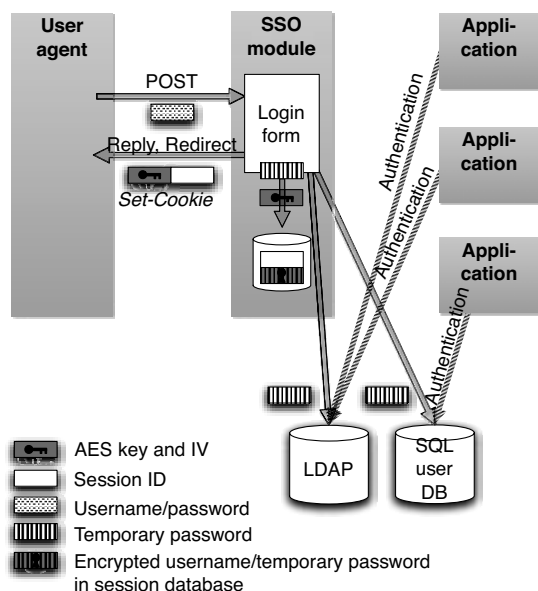


Figure 4: Using temporary passwords for SSO auth

Applications using LDAP BIND may use different LDAP subtrees containing only such temporary elements to benefit from multiple passwords (and to retrieve the user account data values). Applications using LDAP requests directly to read user data and verify it themselves may be reconfigured to use also other attributes as data store. Also, SQL database tables (for SQL authentication) can generally be altered.

The benefits of this system would be that even if an attacker manages to get the temporary password (by getting the AES key, IV, and the session database) for one entry it will be useless after this temporary password expires.

The idea - which is not implemented yet - consists of Polybius creating a random password and pushing it to all authentication sources/databases used by the configured legacy web applications. This random password is then stored in the same way as the “normal” user password would be saved encrypted in the Polybius session database.

6 Performance

6.1 AES decryption with OpenSSL

For our performance tests, the AES implementation built into OpenSSL has been used – OpenSSL is already loaded as a module in Apache. A single thread (single CPU) AES decrypting loop has been run on an AMD Athlon 64 X2 machine several years old, contin-

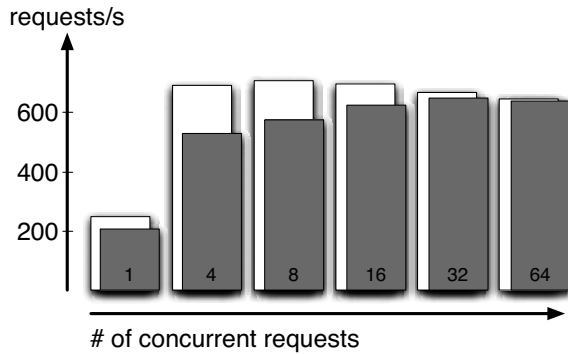


Figure 5: Request rates for proxy passthrough (white) and SSO operation (gray)

uously performing 256 bit key selection followed by AES-CBC decryption on 160 bytes of data. 1 million decryption operations, this setup requires 1.8s, equivalent to $\approx 550,000$ decryptions per second. So AES performance on even an aging machine will not be the bottleneck in this application.

6.2 Latency due to database communications

Latencies occurring when communicating with the database (we used PostgreSQL as the module's storage engine, accessed using prepared statements for efficiency and security). We ran the "ab" Apache benchmark program doing 3000 requests through the Apache Proxy to a static file on another server, also consisting of 160 bytes. To model browsers keeping their sessions to the proxy open and thus save SSL/TLS negotiation time, we used the 'keepalive' feature of "ab". The results are shown in Figure 5, which allow the following interpretation:

1. Performance loss due to the database lookup is almost visible when request concurrency is high, which will be the typical operating point for highly loaded proxies. For weakly loaded proxies, performance will also not be a problem.
2. In the sequential test (concurrency=1), each request is slowed down by 0.8 ms longer (4.8 ms compared to 4.0 ms, resulting in 207 requests/s compared to 249).

Our database log (which I turned on for testing) shows that binding the parameter to the prepared statement plus executing it takes 0.3 ms.⁶ The remaining offset is due to network latency.

This has all been measured on a version of Polybius which has not yet been tuned for optimal performance, as our focus so far has been on functionality and correctness.

⁶ ≈ 0.2 ms to bind to the variable and 0.1ms to execute it and return the result.

7 Conclusions and future work

Recent developments allow the simplification of Single-Sign-On within an organization, as most applications are web-based and even more so in the near future. Server-side Web proxies are commonly used as load-balancing and failover mechanisms, show high performance, and are well known by systems administrators. Even in the unlikely case the SSO proxy should ever fail, the backend applications may still be accessed in the traditional way.

Combining these insights with a secret sharing mechanism allows Polybius to not only create a flexible and easily configurable as well as extensible SSO mechanism, it can also reduce the impact of break-ins or unauthorized disclosure. These combinations make Polybius well-suited for the academic environment with its openness and heterogeneity.

We are working on extending the number of backend protocols supported in Polybius and also believe it will be possible to include Polybius into cross-domain SSO systems such as Shibboleth or certificate-based login mechanisms.

Also, the concept of temporary passwords can be extended to provide (1) flexible user rights delegation or (2) single-use authentication data, enabling the use of a one-time passwords or other restrictions when having to use a public terminal to access one's applications. The beauty of this approach is that, unlike other SSO infrastructure, such mechanisms can be added purely inside the SSO gateway, not a single configuration or source change is needed for the backend applications.

The University of Konstanz is evaluating the usage of Polybius to include their webmail and collaboration applications as well as the central SVN service in a uniform user portal.

References

- [Jas] Jasig. CAS 2 Architecture. <http://www.jasig.org/cas/cas2-architecture>. Accessed 2011-01-11. 2
- [JZB06] Karthik Jaganathan, Larry Zhu, and John Brezak. SPNEGO-based Kerberos and NTLM HTTP Authentication in Microsoft Windows, June 2006. 2
- [NYHR05] Clifford Neuman, Tom Yu, Sam Hartman, and Kenneth Raeburn. The Kerberos Network Authentication Service (V5). IETF RFC 4120, July 2005. 2
- [Sco05] Scott Cantor, editor. Shibboleth Architecture: Protocols and Profiles. <http://shibboleth.internet2.edu/docs/internet2-mace-shibboleth-arch-protocols-200509.pdf>, September 2005. Accessed 2011-01-11. 2

Identity Management

Identitätsmanagement für Hybrid-Cloud-Umgebungen an Hochschulen – Erfahrungen im Münchner Wissenschaftsnetz

Silvia Knittl
Technische Universität München
knittl@tum.de

Wolfgang Hommel
Leibniz-Rechenzentrum
hommel@lrz.de

Abstract: Hybrid-Cloud ist die aktuell gängige Bezeichnung für den Mischbetrieb von lokalen und externen IT-Diensten, die auf Basis einer Kombination aus physischer und virtueller Hardware erbracht werden. Die aktuellen Entwicklungen im Cloud-Computing bringen eine Reihe interessanter Management- und Administrationswerkzeuge hervor, deren Einsatz auch für Hochschulen und deren Rechenzentren attraktiv ist, selbst wenn dort nicht jeder neue Hype mit offenen Armen empfangen wird. Die neuen Formen der Dienstleistung, die sich dank der neuen Werkzeuge inzwischen auch effizient in die Praxis umsetzen lassen, bringen aber auch eine Vielzahl neuer Herausforderungen mit sich, die sich wiederum auf die bereits vorhandene IT-Infrastruktur auswirken. In diesem Beitrag beleuchten wir das Zusammenspiel zwischen Hybrid-Cloud-Umgebungen und dem Identity & Access Management, also einer der tragenden Säulen des organisationsweiten Sicherheitsmanagements, im Hochschulumfeld auf Basis unserer praktischen Erfahrungen im Münchner Wissenschaftsnetz.

1 Motivation

Wie in privatwirtschaftlichen Unternehmen hat der Betrieb von IT-Systemen auch in Hochschulumgebungen – von Forschungsvorhaben zunächst abgesehen – keinen Selbstzweck, sondern dient der gezielten Unterstützung der Geschäfts- bzw. Hochschulprozesse. Heute könnte keine Hochschule ihren Aufgaben in einer für Mitarbeiter und Studierende attraktiven Weise mehr nachkommen, ohne Anwendungen, wie etwa Personal- und Studentenverwaltungssoftware, E-Mail, Web- und Fileserver, einzusetzen. Anders als in den meisten Unternehmen erfolgt der IT-Betrieb bislang aber nicht stark organisationsintern zentralisiert bzw. mit einer zentralisierten Koordination möglicher Outsourcing-Bereiche. Vielmehr stellt eine Verteilung der IT-Ressourcen, beispielsweise zwar mit einem Schwerpunkt im Hochschulrechenzentrum, aber mit deutlich erkennbaren Mengen ergänzender Systeme in der Hochschulverwaltung und den einzelnen Fakultäten und Fachbereichen, ein nach wie vor gängiges Bild vieler deutscher Hochschulen dar.

Ein Wandel dieser Form des IT-Betriebs wird dabei einerseits durch Rezentralisierungsbemühungen angestrebt, der den in den 1990er Jahren vorherrschenden IT-Dezentralisierungsdrang eindämmen, Wildwuchs verhindern, Kompetenzen bündeln und die Gesamtkosten beim Betrieb der Hochschul-IT-Infrastruktur reduzieren soll, aber aufgrund

befürchteter resultierender Einschränkungen für Forschung und Lehre oft auch auf Widerstand stößt. Andererseits ändern sich durch die natürliche Fluktuation an Hochschulen und die massiv zunehmende Anzahl so genannter *Digital Natives* auch die Anforderungen an die IT-Dienste an einer Hochschule, da Personen, die mit IT-Systemen und IT-Diensten aufgewachsen sind und sich auch privat intensiv damit auseinandersetzen, ein offenkundig meist grundlegend anderes Verständnis von IT mitbringen als beispielsweise die Angehörigen geisteswissenschaftlicher Fachrichtungen vor 20 Jahren. Gerade diese Zielgruppe, die im weitesten Sinne als Cloud-Services bezeichnbare IT-Dienste aus ihrem privaten Umfeld kennt, wünscht sich auch an der Hochschule vergleichbare Dienstleistungen und stellt dabei die klassische Prämisse, möglichst viel Hardware im eigenen Bürotrakt unterzubringen, in den Hintergrund [BDNP10].

Die Beurteilung von Cloud-Computing an deutschen Hochschulen gestaltet sich schwierig: Während sich insbesondere Informatik und Wirtschaftswissenschaften durchaus überwiegend euphorisch mit dem neuen Forschungsgebiet auseinandersetzen und beispielsweise seine ökonomischen, betrieblichen und sicherheitsspezifischen Aspekte durchleuchten, sehen andere darin alten Wein in neuen Schläuchen oder einen durch die Marketingabteilungen von Herstellern hochgehaltenen Hype und fragen sich, wo der Unterschied zum in der Wissenschaft längst etablierten Grid Computing oder konkrete Anwendungsfälle an Hochschulen liegen sollen (vgl. [FZRL09]). Allen Begriffsunklarheiten und Definitionsversuchen, die das Cloud-Computing immer noch prägen, zum Trotz zeigen sich jedoch zwei Eigenschaften, die auch für Hochschulen interessant sind: Zum einen kommt dem Einsatz von System- und Netzvirtualisierungstechniken eine tragende Rolle zu, die unabhängig vom Cloud-Hype verstärkt in den letzten 3–5 Jahren auch in Hochschulrechenzentren Einzug gehalten haben. Zum anderen bietet Cloud-Computing mit der Kategorisierung in Infrastructure, Platform und Software as a Service (IaaS, PaaS und SaaS) eine auch auf klassische Hochschul-IT-Dienste anwendbare Strukturierung, deren praktische Bedeutung mit der zunehmenden Verfügbarkeit entsprechender Management- und Administrationswerkzeuge zunimmt.

In diesem Beitrag betrachten wir das Cloud-Computing weder direkt noch in seinem vollen Umfang. Vielmehr konzentrieren wir uns nach einem kurzen Überblick über eine ausgewählte, sich im Produktivbetrieb befindende „Hochschul-Cloud“ auf so genannte *hybrid clouds*. Eine Hybrid-Cloud liegt nach aktueller allgemeiner Auffassung dann vor, wenn eine Kombination verschiedener Cloud-Modelle (public, private, community; häufig auf virtueller Hardware basierend) mit einer traditionellen IT-Umgebung (auf physischer Hardware basierend) vorliegt. Die Begriffe *public*, *private* und *community* beziehen sich hierbei auf die Betreibermodelle bzw. Eigentumsverhältnisse (siehe [MG09]). Durch eine geeignete Kombination sollen a) die Vorteile aller Varianten entsprechend zum Tragen kommen, b) mit zusätzlichen Cloud-Ressourcen Lastspitzen (z. B. zu Semesterbeginn) abgedeckt werden und c) sanfte Migrationen bzw. schnelle Integrationsmöglichkeiten geboten werden (vgl. [KP10]). Neue technische Möglichkeiten bringen im Allgemeinen aber auch neue Risiken mit sich, und eine Hybrid-Cloud bildet in dieser Hinsicht leider keine Ausnahme (vgl. [ENI09]). Von den vielen, bislang meist noch nicht ausreichend erforschten und nur unzureichend praktisch gelösten IT-Security-Problemen, die mit dem Cloud-Computing einhergehen, konzentrieren wir uns bei den Hybrid-Clouds im Sinne eines den Daten-

schutz betonenden *Privacy-by-Design*-Ansatzes (vgl. [Cav10]) auf das Identity & Access Management (I&AM), wie es sich typischerweise zur Benutzer- und Berechtigungsverwaltung, beispielsweise mit LDAP-Servern und Konnektoren zu den Quellsystemen im Campus Management implementiert, an Hochschulen findet.

Als Beispiel ziehen wir dafür die Zusammenarbeit zwischen der Technischen Universität München (TUM) und dem Leibniz-Rechenzentrum (LRZ) im Münchner Wissenschaftsnetz (MWN) heran. Das LRZ ist der zentrale IT-Dienstleister aller Münchner Wissenschaftseinrichtungen (siehe [LRZ10]) und hat sein I&AM-System eng mit den entsprechenden Systemen insbesondere der Ludwig-Maximilians-Universität (LMU) und der TUM gekoppelt, um den Hochschulangehörigen einen möglichst benutzerfreundlichen Zugang zu den von ihnen benötigten IT-Diensten zu ermöglichen. Die im Rahmen des 2009 abgeschlossenen, DFG-geförderten Projekts IntegraTUM mit dem Ziel der Übertragbarkeit ihrer Bestandteile auf andere Hochschulen entwickelte I&AM-Infrastruktur [BEH⁺10] wird TUM-weit genutzt, regelt die Nutzung von LRZ-Diensten durch TUM-Angehörige und dient auch maßgeblich der Integration innerhalb der TUM neu aufgebauter IT-Dienste. Abbildung 1 zeigt einen vereinfachten Ausschnitt der IT-Dienstleistungslandschaft der TUM mit eigenen IT-Ressourcen in Kombination mit Diensten von public Cloud-Providern (Wikispaces in der Abbildung) und dem LRZ als IT-Dienstleister, wobei die einzelnen IT-Dienste den drei Kategorien IaaS, PaaS und SaaS zugeordnet sind. Die Motivation für und die Details dieser Zuordnung finden sich in unserer früheren Arbeit [KL10]; Ansätze zum hochschulspezifischen IT-Risikomanagement haben wir im Kontext des MWN in [KH10] vorgestellt.

Im nächsten Abschnitt erörtern wir anhand dieses Beispiels die Auswirkungen von Hybrid-Clouds auf Hochschul-I&AM-Systeme und schließen mit einer Zusammenfassung der Ergebnisse in Abschnitt 3.

2 Hochschul-Identity-Management und der Betrieb von Hybrid-Clouds

I&AM-Systeme lassen sich auch bei ihrem Einsatz an Hochschulen bezüglich ihrer Gesamtarchitektur in drei Bereiche untergliedern, die auch in Abbildung 2 am Beispiel von TUM und LRZ dargestellt sind:

- Autoritative Quellsysteme sind eng mit den Geschäftsprozessen verknüpft und liefern dem Kern des I&AM-Systems möglichst alle zu berücksichtigenden Benutzer. An Hochschulen stellen traditionell die Personal- und Studierendenverwaltungssysteme die wichtigsten Datenquellen dar und werden z. B. durch Gästeverwaltungsmechanismen (für Konferenzteilnehmer, Angebote für Schüler u. ähnl.) ergänzt. Die Betrachtung einzelner Datenquellen ist in den vergangenen Jahren dem Paradigma der Campus-Management-Systeme (CaMS) gewichen, die beispielsweise auch sicherstellen, dass ein Studierender, der gleichzeitig auch Mitarbeiter ist (z. B. als Hilfskraft oder im Rahmen eines Promotionsstudiengangs), nur einmal und nicht mehrfach als digitale Identität erfasst wird.

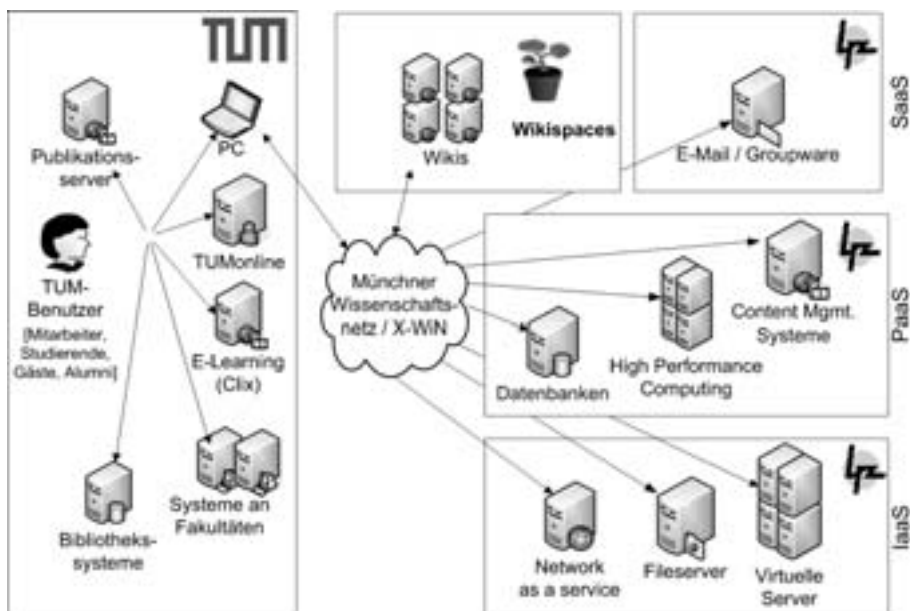


Abbildung 1: Die IT-Dienstlandschaft der TUM im Zusammenspiel mit dem LRZ

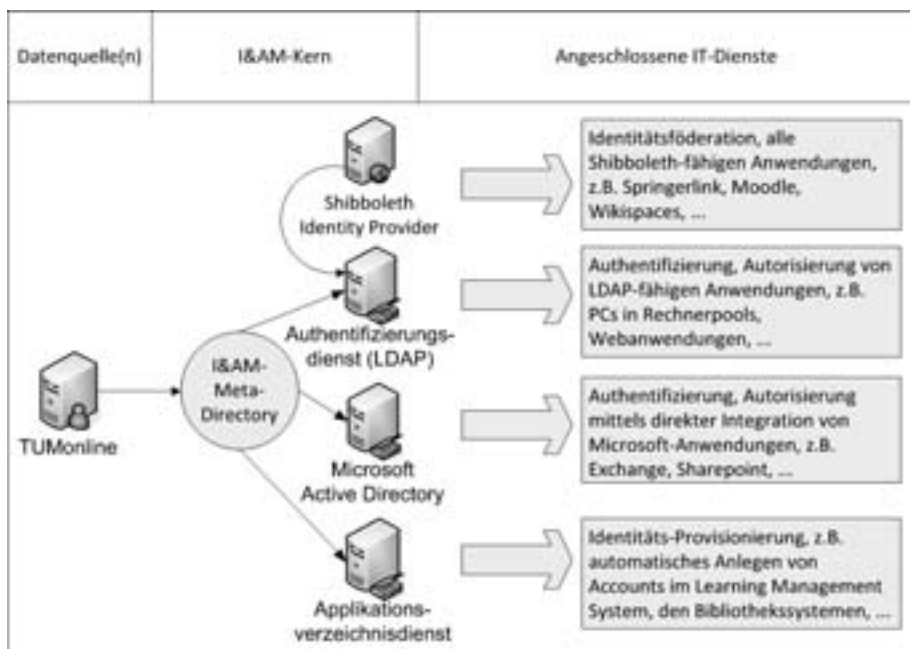


Abbildung 2: Die drei Bereiche Quellen, Kern und Dienste in I&AM-Architekturen

- Der Kern des I&AM-Systems besteht aus einem zentralen Datenbestand, der meist in Form von LDAP-Servern bereitgestellt wird, und auf diesem aufbauenden Funktionen, so dass beispielsweise aus der organisatorischen Zugehörigkeit eines Benutzers automatisch Default-Berechtigungen abgeleitet werden können. Die Grenzen zwischen den von CaMS und den von I&AM-Self-Service-Portalen erbrachten Möglichkeiten zur manuellen Steuerung individueller Berechtigungen sind derzeit fließend: Während beispielsweise das Ändern von Passwörtern durch Benutzer klassisch dem I&AM-System zuzuordnen ist, fungieren CaMS zunehmend als zentrale, webbasierte Schnittstelle auch zu den einzelnen Benutzern und reichen entsprechende Datenänderungen an das I&AM-System durch.
- Die einzelnen IT-Dienste und IT-Systeme nutzen das I&AM-System als technische Komponente zur Authentifizierung und Autorisierung ihrer Benutzer. Während LDAP-fähige IT-Dienste auf einen lokalen Benutzerdatenbestand ganz verzichten oder lediglich ergänzende Informationen lokal vorhalten müssen, werden die anderen IT-Dienste vom I&AM-System über so genannte Konnektoren mit allen relevanten Benutzerdaten versorgt; der entsprechende Vorgang wird als (User) Provisioning bezeichnet.

Für rein hochschulintern genutzte IT-Dienste ist das Thema Benutzerverwaltung mit einer Anbindung an das I&AM-System erledigt: Die Integration möglichst vieler Hochschul-IT-Dienste und der dafür pro IT-Dienst erforderliche Aufwand sind weit verbreitete (aber nicht die einzigen) Beurteilungskriterien für I&AM-Systeme. Sollen darüber hinaus auch Externe auf die IT-Dienste Zugriff erhalten, werden Mechanismen des Federated Identity Management angewandt, beispielsweise durch eine Integration in die Authentifizierungs- und Autorisierungsinfrastruktur des DFN-Vereins (DFN-AAI); in diesen Fällen steht, beispielsweise durch den Einsatz der Software Shibboleth, nicht mehr nur der lokale LDAP-Server zur Benutzerverwaltung zur Verfügung, sondern es können beispielsweise auch ausgewählte Daten externer Studierender von deren jeweiliger Heimathochschule abgefragt werden. Eine praktische Anwendung erfolgt beispielsweise bei Learning Management Systemen über das DFN-AAI E-Learning-Profil [DGH⁺08]: Studierende einer Hochschule können E-Learning-Kurse an anderen Hochschulen belegen und erhalten dafür auch Leistungsnachweise; dieser Prozess ist beispielsweise an der Virtuellen Hochschule Bayern (vhb) seit längerem in Betrieb und wird derzeit nach und nach in die DFN-AAI integriert.

Das Identity-Management an Hochschulen dient aber in den meisten Fällen vorrangig der Bewältigung eines Massenproblems: Viele Benutzer und zahlreiche angeschlossene Systeme müssen möglichst weitgehend automatisiert samt ihren Berechtigungen verwaltet werden. Während das I&AM-System also durchaus dem E-Mail-System gegenüber als autoritative Datenquelle bezüglich der E-Mail-Adressen der eigenen Benutzer fungiert und dem Learning Management System anhand von Informationen über den vom Studierenden belegten Studiengang liefert, um den Zugang zu ausgewählten Kursen freizuschalten, wurden systemadministrative Berechtigungen – beispielsweise, wer auf welchen Linux-Servern privilegierten Zugriff hat – bislang nur unzureichend in der Kette CaMS–I&AM–Dienst betrachtet. So genannte Privileged Account Management Systeme, die sich auf die Verwaltung solcher Kennungen mit weitreichenden Befugnissen und Maßnahmen gegen ihren

Missbrauch befassen, haben zwar längst einen festen Platz beispielsweise im militärischen Bereich und im Bankenwesen, dringen unter dem Einfluss internationaler gesetzlicher Auflagen u. a. in multinationale börsennotierte Unternehmen vor, spielten an Hochschulen aber bislang keine große Rolle, da die Anzahl der Administratoren überschaubar und der mutwillige Missbrauch damit einhergehender Berechtigungen selten war.

Mit dem Betrieb von Hybrid-Clouds geht nun jedoch einher, dass durch das einfache Bereitstellen virtueller Maschinen nicht nur die Gesamtzahl zu verwaltender Systeme insgesamt zunimmt, sondern sich auch die Anzahl mit erweiterten Berechtigungen ausstattender Benutzer erhöht. In Kombination mit der in der Praxis leidlich hohen Dynamik, die das einfache Hinzuschalten und Wegnehmen virtueller Maschinen ermöglicht, ist eine wie bislang übliche, manuelle Verwaltung von Administrator-Accounts nicht mehr sinnvoll. Beispielsweise plant das LRZ derzeit die Inbetriebnahme eines als VM-Shop bezeichneten Webportals, mit dem von Lehrstuhlinhabern berechnete Hochschulmitarbeiter virtuelle Maschinen in einem vorab vereinbarten Maximalumfang online konfigurieren und weitestgehend automatisiert in Betrieb nehmen können. Mit der Abwicklung der Bestellung ist verbunden, dass der Anwender eine Reihe zusätzlicher Berechtigungen erhält: Neben den systemadministrativen Berechtigungen auf der neuen virtuellen Maschine muss er beispielsweise für die Nutzung der VMware-Managementwerkzeuge, die für das Ein- und Ausschalten der virtuellen Maschine benötigt werden und weitere Funktionen wie Snapshots zur Verfügung stellen, freigeschaltet werden, was aufgrund der Integration der VMware-Umgebung in die Microsoft-Windows-Umgebung entsprechende Berechtigungen im Active Directory voraussetzt (vgl. Abbildung 3). Somit ergeben sich bereits für relativ einfache Anwendungsfälle komplexe Abhängigkeiten, die das I&AM-System aufgrund der hohen Dynamik vollständig automatisieren muss, und deren Umsetzung mangels einfacher manueller Kontrollmethoden auch konsequent automatisch auf Konsistenz und mögliche Fehler (wie fälschlicherweise nicht wieder entzogene Berechtigungen) geprüft werden muss.

Die rein anwendungsfallgetriebene Erweiterung von Hochschul-I&AM-Systemen wie im Beispiel für die Bereitstellung kundenspezifischer virtueller Maschinen würde langfristig aber nicht skalieren, da für jeden neuen Typ von Cloud-Dienst, der angeboten werden soll, spezifische Erweiterungsarbeiten erforderlich wären. Auch im Hinblick auf die hochschulrechenzentrums-interne Umsetzung des Hybrid-Cloud-Ansatzes bietet sich deshalb eine Orientierung an den von der Cloud Security Alliance (CSA) im April 2010 veröffentlichten Leitlinien für das I&AM [Clo10] an. Dementsprechend sind vier funktionale Bereiche zu betrachten:

1. Provisioning und Deprovisioning digitaler Identitäten: Das bislang übliche reine Anlegen und Löschen neuer digitaler Identitäten in einem System, zu dessen Benutzung der Anwender berechtigt wurde bzw. dem die Berechtigung entzogen wurde, alleine reicht nicht mehr aus. Vielmehr untergliedert sich die Umsetzung einer Berechtigungsänderung in einzelne Schritte (von der CSA als *multi-stage setup* bezeichnet), mit denen auch Einträge und Zuordnungen z. B. in den Managementsystemen vorgenommen werden. Dabei muss ggf. auch auf die richtige Reihenfolge geachtet werden, was bei vielen Implementierungen, die Provisioning-Aktivitäten hochgradig parallelisieren, zur Erforderlichkeit tiefer Eingriffe führen kann.

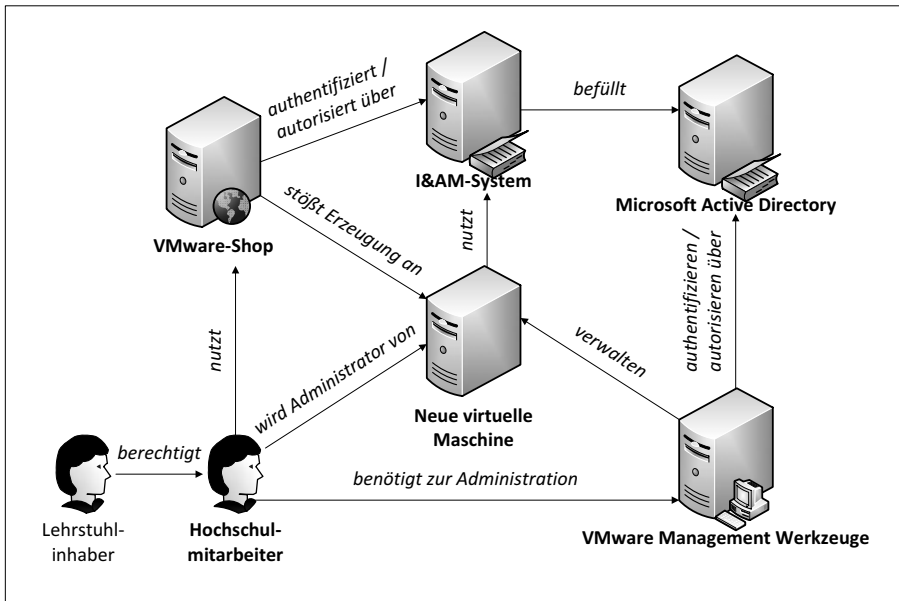


Abbildung 3: Abhängigkeiten bei der Bereitstellung virtueller Maschinen im MWN

2. Authentifizierung, Single Sign-On und Föderation: Für bereitgestellte Cloud-Services muss z. B. das Thema Passwort-Management überdacht werden; beispielsweise sollte vermieden werden, dass neue virtuelle Maschinen entweder alle dasselbe Startpasswort für den root-Account verwenden oder dass das benutzerspezifische Passwort daran gekoppelt wird (wie es bei vielen anderen Diensten üblich ist, damit der Benutzer ein gemeinsames Passwort für alle seine persönlichen Hochschul-IT-Dienste hat). Die dynamisch bereitgestellten Dienste sollten zudem bezüglich ihrer Benutzerverwaltung so vorkonfiguriert werden, dass sie vorhandene I&AM-Systeme bzw. Föderationsschnittstellen nutzen können, ohne dass ein lokaler Benutzerdatenbestand aufgebaut werden muss.
3. Zugriffskontrolle und Benutzerprofilmanagement: Mit dem Provisioning wird auch die Zugriffsüberwachung auf die Cloud-Ressourcen komplexer. Beispielsweise müssen auch Monitoringsysteme an die Managementwerkzeuge für virtuelle Maschinen gekoppelt werden, um entscheiden zu können, ob eine virtuelle Maschine ausgefallen ist oder vom Benutzer bewusst abgeschaltet wurde. Analog dazu können die Anwender von Cloud-Diensten anderen Anwendern selektiv Zugriff gewähren, um beispielsweise Mashups implementieren zu können, was über das I&AM-System und die Benutzerprofile geeignet abgebildet werden muss. Ebenso müssen Protokolle und Audit-Logs nicht mehr nur intern ausgewertet, sondern auch den entsprechenden Anwendern für ihren jeweiligen Bereich zur Verfügung gestellt werden.
4. Compliance: Gesetzliche Auflagen spielen nicht nur dann eine Rolle, wenn externe

Anbieter mit der Bereitstellung von Cloud-Ressourcen beauftragt werden, worauf wir hier nicht näher eingehen. Vielmehr müssen beispielsweise adäquate Maßnahmen zum Datenschutz auch und gerade dann umgesetzt werden, wenn neue Dienste dynamisch im Auftrag eines Benutzers instanziiert werden. Hierzu gehört einerseits, die unnötig breite Streuung personenbezogener Daten auf viele Systeme zu vermeiden, und zum anderen die Implementierung von Mechanismen, um anfallende personenbezogene Daten wie Zugriffsprotokolle und detaillierte Accountinginformationen innerhalb eines angemessenen Zeitraums auch wieder zu löschen.

Für die in Hochschul-I&AM-Systemen typischerweise enthaltenen Komponenten ergeben sich somit folgende Auswirkungen:

- Konnektoren und Verzeichnisdienste mit dienstspezifischen Datenschemata: Die gezielte Versorgung einzelner IT-Dienste mit Benutzerdaten im von ihnen benötigten Format erfordert aufwendige Implementierungsarbeiten, z. B. zur Datentransformation zwischen dem I&AM-Format und anwendungsspezifischen Benutzerprofilformaten. Die Anbindung über Konnektoren eignet sich auch unabhängig von diesem Bereitstellungsaufwand nur für längerfristig betriebene, einzelne IT-Dienste. Für die Integration einer Vielzahl dynamisch bereitgestellter Dienste muss auf eine der anderen Varianten ausgewichen werden.
- LDAP-basierte Authentifizierungs- und Autorisierungsdienste: Durch ihre im allgemeinen hohe Performance bzw. Skalierbarkeit und die clientseitig einfache Konfiguration sind LDAP-Server auch für die Anbindung der dynamisch über Cloud-Methoden eingerichteten Dienste attraktiv. Sie eignen sich insbesondere für IT-Dienste, bei denen sich der an der Hochschule registrierte Benutzer explizit authentifizieren soll. Falls hingegen Eigenschaften wie Single Sign-On oder die Unterstützung externer Benutzer über Föderationen benötigt werden, sollte auf eine der beiden nachfolgenden Varianten gesetzt werden.
- Active-Directory-Integration: Zum zentralen Management von Windows-Servern und Desktop-PCs bzw. Notebooks mit Microsoft-Betriebssystem kommt an sehr vielen Hochschulen bereits ein Microsoft Active Directory zum Einsatz; es ist wie oben erläutert auch eine Voraussetzung für den Betrieb einer VMware-basierten Virtualisierungsinfrastruktur und bildet die Basis für diverse weitere Microsoft- und Windowsdienste wie Exchange und Sharepoint. Durch seine LDAP-Schnittstelle kann es jedoch auch für die Anbindung von Linux-Maschinen genutzt werden. Mit dem Einsatz von Active Directory ist u. a. mittels seiner Kerberos-Funktionalität die Möglichkeit zum Single Sign-On, also die Nutzung aller ins Active Directory integrierten Dienste ohne pro Dienst zu wiederholende Passworteingabe, gegeben. Sie bietet sich deshalb insbesondere zur nahtlosen Integration der mit den Cloud-Diensten assoziierten Managementwerkzeuge an.
- Shibboleth (oder andere Software zur Teilnahme an hochschulübergreifenden Föderationen): Die Anpassung noch nicht föderationsfähiger Software an Shibboleth ist nach wie vor mit einem nicht zu vernachlässigenden Einarbeitungsaufwand für

die Entwickler bzw. Administratoren verbunden. Allerdings steigt die Anzahl der auch im Hochschulumfeld populären Webanwendungen, die an Shibboleth angepasst wurden, weiterhin stetig an. Neben dem Single Sign-On wird mit dieser Variante auch die selektive und aus Anwendungssicht transparente Benutzung durch externe Anwender unterstützt. Da I&AM-seitig keine Anpassungen erforderlich werden, skaliert dieser Lösungsansatz auch, wenn eine Vielzahl neuer Dienste dynamisch mittels Cloud-Ressourcen aufgebaut wird.

Die sich für den Betrieb Cloud-basierter Dienste abzeichnende Abkehr von konnektoren-basiertem Provisioning und die verstärkte Integration in Active-Directory-Infrastrukturen erfordert somit ein partielles Umdenken und tiefere Eingriffe in die bereits implementierten I&AM-Systeme.

3 Zusammenfassung

In diesem Beitrag haben wir zunächst motiviert, dass es im Umfeld des Cloud-Computing durchaus Entwicklungen gibt, die auch für Hochschulen und deren Rechenzentren interessant sind. Hybrid-Clouds, die eigene IT-Dienste mit Cloud-Services gezielt miteinander kombinieren, eignen sich unter anderem dazu, die Vorteile beider Betriebsformen zu kombinieren, Lastspitzen abzufedern und Migrationsszenarien umzusetzen. Die Bereitstellung virtueller Maschinen, eine der grundlegenden Technologien des Cloud-Computings, ist für viele *Digital Natives* darüber hinaus schon zum Normalfall im privaten Umfeld geworden, so dass entsprechende Anforderungen auch an Hochschul-IT-Infrastrukturen laut werden.

Die neue Technologie hat aber auch Auswirkungen auf schon vorhandene Konzepte und Architekturen, beispielsweise auf das Identity & Access Management. Am Beispiel des Münchner Wissenschaftsnetzes haben wir aufgezeigt, mit welchen Zielsetzungen und Realisierungsoptionen I&AM-Systeme an Hochschulen in den letzten Jahren aufgebaut wurden und welche Änderungen erforderlich werden, um die Anforderungen, die sich aus dem Einsatz von Cloud-Technologien ergeben, erfüllen zu können. Dabei haben wir uns am I&AM-Leitfaden der Cloud Security Alliance orientiert und seine Konzepte auf die Hochschul-I&AM-Landschaft übertragen.

Diese Untersuchungen haben gezeigt, dass bei I&AM-Integrationskonzepten und den technischen Infrastrukturen zu ihrer Umsetzung tiefere Eingriffe erforderlich sind, um beispielsweise ein Privileged Account Management und gegenseitige Abhängigkeiten zwischen Berechtigungen adäquat berücksichtigen zu können. Während sich bereits an vielen Hochschulen begonnene Aktivitäten zur Anpassung von Diensten an Shibboleth auch im Kontext von Cloud-Diensten bezahlt machen, stoßen traditionelle I&AM-Konzepte wie der Einsatz von Konnektoren zum Provisioning bei der Cloud-Dynamik an ihre Grenzen.

Die Umsetzung der neuen Konzepte wird im Münchner Wissenschaftsnetz derzeit mit dem Aufbau eines VM-Shops erprobt, der autorisierten Hochschulmitarbeitern u. a. der LMU und TUM ermöglichen soll, schnell und weitestgehend automatisiert neue virtuelle Maschinen bereitzustellen, um auf dieser Basis eigene Hochschul-IT-Dienste anzubieten, die

nahtlos in die vorhandene I&AM-Infrastruktur und die Nutzung der DFN-AAI-Föderation integriert werden können.

Danksagung:

Die Autoren danken den Mitgliedern des Munich Network Management Teams (MNM-Team) für wertvolle Kommentare und Anregungen zu früheren Versionen dieses Beitrags. Das MNM-Team ist eine Forschungsgruppe der Ludwig-Maximilians-Universität, der Technischen Universität München, der Universität der Bundeswehr München und des Leibniz-Rechenzentrums der Bayerischen Akademie der Wissenschaften unter der Leitung von Prof. Dr. Dieter Kranzlmüller und Prof. Dr. Heinz-Gerd Hegering. Siehe <http://www.mnm-team.org/>.

Literatur

- [BDNP10] Rob Bristow, Ted Dodds, Richard Northam und Leo Plugge. Cloud Computing and the Power to Choose. *EDUCAUSE Review*, 45(3), Mai/Juni 2010. Verfügbar online unter <http://net.educause.edu/ir/library/pdf/ERM1030.pdf>.
- [BEH⁺10] Latifa Boursas, Ralf Ebner, Wolfgang Hommel, Silvia Knittl und Daniel Pluta. IntegraTUM Teilprojekt Verzeichnisdienst: Identity & Access Management als technisches Rückgrat der Hochschul-IuK-Infrastruktur. In Arndt Bode und Rolf Borgeest, Herausgeber, *Informationsmanagement in Hochschulen*. Springer Berlin Heidelberg, 2010.
- [Cav10] Ann Cavoukian. Modelling Cloud Computing Architecture Without Compromising Privacy: A Privacy by Design Approach. Verfügbar online unter <http://www.ipc.on.ca/images/Resources/pbd-NEC-cloud.pdf>, Mai 2010.
- [Clo10] Cloud Security Alliance. Domain 12: Guidance for Identity & Access Management - V2.1. Verfügbar online, April 2010.
- [DGH⁺08] Jörg Deutschmann, Peter Gietz, Wolfgang Hommel, Renate Schroeder, Jens Schwendel und Tobias Thelen. DFN-AAI E-Learning-Profil: Technische und organisatorische Voraussetzungen, Attribute für den Bereich E-Learning. <http://www.aai.dfn.de/der-dienst/attribute/>, 2008.
- [ENI09] European Network and Information Security Agency ENISA. Cloud Computing - Benefits, risks and recommendations for information security. Verfügbar online unter <http://www.enisa.europa.eu/act/rm/files/deliverables/cloud-computing-risk-assessment>, Zugriff am 12.01.2010, November 2009.
- [FZRL09] I. Foster, Y. Zhao, I. Raicu und S. Lu. Cloud Computing and Grid Computing 360-Degree Compared. *ArXiv e-prints*, 901, December 2009.
- [KH10] Silvia Knittl und Wolfgang Hommel. Umgang mit Risiken für IT-Dienste im Hochschulumfeld am Beispiel des Münchner Wissenschaftsnetzes. In *INFORMATIK 2010, Band 2*, 2010.
- [KL10] Silvia Knittl und Albert Lauchner. Hybrid-Cloud an der Technischen Universität München – Auswirkungen auf das IT-Management. In *INFORMATIK 2010, Band 1*, 2010.
- [KP10] Silvia Knittl und Hans Pongratz. Application Integration Methods For Learning Management Systems. In IADIS, Herausgeber, *International Conference e-Learning 2010*, Freiburg, Deutschland, Juli 2010.
- [LRZ10] LRZ. Dienstleistungskatalog. Verfügbar online unter <http://www.lrz.de/wir/regelwerk/dienstleistungskatalog.pdf>, September 2010.
- [MG09] Peter Mell und Tim Grance. The NIST Definition of Cloud Computing. Technical report, National Institute of Standards and Technology, Information Technology Laboratory, Juli 2009.

Benutzerzentrierung und Föderation: Wie kann ein Benutzer die Kontrolle über seine Identitätsdaten bewahren?

Holger Kuehner, Thorsten Hoellrigl, Hannes Hartenstein
Karlsruher Institut für Technologie, Steinbuch Centre for Computing
{holger.kuehner|hoellrigl|hartenstein}@kit.edu

Abstract:

Viele Dienste benötigen zur Dienstbringung, insbesondere zur Absicherung des Dienstzugriffs, identitätsbezogene Informationen. Da sich identitätsbezogene Informationen über die Zeit ändern, muss deren Aktualität zur Vermeidung einer fehlerhaften Dienstbringung das Ziel verteilter IdM-Systeme sein. Angesichts der Vielfalt der Dienstlandschaft steht der Benutzer vor der Herausforderung, die Weitergabe seiner Identitätsinformationen zu kontrollieren sowie zu überblicken, welcher Dienst welche Informationen vorhält. Das benutzerzentrierte föderative Identitätsmanagement versucht dieser Herausforderung zu begegnen, indem es den Benutzer bei der Verwaltung seiner Identitätsinformationen in den Mittelpunkt stellt. Aktuelle Ansätze in diesem Bereich weisen jedoch Schwachstellen auf, vor allem durch die Notwendigkeit einer *manuellen Zustimmung* zu jeder Weitergabe identitätsbezogener Informationen durch den Benutzer. Folglich können auftretende Änderungen nur verteilt werden, wenn der Benutzer online ist, was wiederum zu einer unzureichenden Aktualität dieser Informationen führen kann. In den Arbeiten [HKDH10, HKDH11] wurde ein Ansatz namens *User-Controlled Automated Identity Delegation* (UCAID) präsentiert, mittels welchem eine Verbesserung der Aktualität identitätsbezogener Informationen unter Berücksichtigung des Datenschutzes und der Datensicherheit erreicht werden kann. Die Entwicklung und Darstellung von Ansätzen zur Verwaltung identitätsbezogener Informationen orientiert sich derzeit noch stark an einzelnen Aspekten mit teilweise unklarer Bedeutung wie bspw. „Benutzerfreundlichkeit“. Hierunter kann das Verständnis des Funktionsumfangs der einzelnen Ansätze leiden, was wiederum die Vergleichbarkeit der Ansätze einschränkt. Um einen Vergleich verschiedener Ansätze erreichen zu können, werden in dieser Arbeit zunächst relevante Dimensionen aufgezeigt, bspw. die Art der Initiierung des Informationsaustauschs, und anhand von UCAID dargelegt, wie sich mit Hilfe dieser Dimensionen Ansätze zur verteilten Verwaltung identitätsbezogener Informationen systematisch entwerfen und miteinander vergleichen lassen. Des Weiteren wird hinsichtlich der Initiierung des Informationsaustauschs ein Vergleich Push- und Pull-basierter Ansätze basierend auf einer prototypischen Implementierung des Ansatzes vorgenommen.

1 Einleitung

Ein Dienstanbieter, welcher Benutzern einen gesicherten Zugriff auf seine Dienste bereitstellen möchte, muss hierfür fundamentale Prozesse wie die Identifizierung, die Authentifizierung und die Autorisierung umsetzen. Die Basis zur Realisierung dieser Prozesse

bilden identitätsbezogene Informationen, wie bspw. die Anschrift, eine Organisationszugehörigkeit oder auch der Nachname des Benutzers. Da diese Informationen Änderungen unterliegen, gilt es zur Vermeidung einer fehlerhaften Dienstleistung oder einer unbeabsichtigten Dienstverweigerung, deren Aktualität sicherzustellen. Diese Herausforderung gewinnt heutzutage insbesondere dadurch an Gewicht, dass die Anzahl der durch einen Benutzer verwendeten Dienste stetig steigt.

Das benutzerzentrierte föderative Identitätsmanagement (bFIM) versucht dieser Herausforderung zu begegnen, indem es den Benutzer bei der Verwaltung seiner Identitätsinformationen in den Mittelpunkt stellt. Aktuelle Ansätze in diesem Bereich machen es jedoch notwendig, dass ein Benutzer jeder Weitergabe identitätsbezogener Informationen manuell zustimmt. Dies kann in einer für manche Anwendungsszenarien unzureichenden Aktualität identitätsbezogener Informationen resultieren.

In [HKDH10, HKDH11] wurde ein Ansatz namens *User-Controlled Automated Identity Delegation (UCAID)* vorgeschlagen, mittels welchem eine verbesserte Aktualität identitätsbezogener Informationen erreicht werden kann. Dieselbe Zielsetzung kann auch über Ansätze wie bspw. UMA [MMCV10] erreicht werden, welches auf dem Autorisierungsmechanismus OAuth [OAU10] basiert. Weitestgehend ungeklärt blieb jedoch die Fragestellung, wie sich diese und weitere Ansätze zur verteilten Verwaltung identitätsbezogener Informationen systematisch entwickeln und miteinander vergleichen lassen.

Die Realisierung eines Ansatzes zur verteilten Verwaltung von Identitätsinformationen befindet sich hierbei in einem Spannungsfeld, welches im Wesentlichen durch den *Benutzer* auf der einen Seite und die *Dienstföderation* auf der anderen Seite aufgespannt wird. Die Dienstföderation bilden hierbei Identitätsanbieter (IdPs), welche identitätsbezogene Informationen vorhalten und weitergeben, sowie Dienstanbieter (SPs), welche diese Informationen konsumieren. Eine Folge dieses Spannungsfeldes ist es, dass zum einen die Verwaltung von Identitätsinformationen möglichst benutzerzentriert realisiert werden sollte, d.h. möglichst benutzerfreundlich und unter Berücksichtigung des Datenschutzes. Zum anderen müssen Anforderungen der Dienstanbieter, welche teilweise konträr zu den Anforderungen des Benutzers sein können, ebenfalls adressiert werden. Eine Betrachtung dieser *Dimensionen* in Kombination mit der Berücksichtigung des Spannungsfeldes zwischen *Benutzerzentrierung und Föderation* sowie den hieraus resultierenden Freiheitsgraden hinsichtlich der Entwicklung von Ansätzen zur Verwaltung identitätsbezogener Informationen in verteilten Systemen soll hierbei den Mehrwert liefern, dass Ansätze *systematisch* entwickelt und bewertet werden können. Folgendes trägt die vorliegende Arbeit dazu bei:

- (i) Eine Darlegung relevanter Dimensionen, welche bei der Verwaltung verteilter Identitätsinformationen berücksichtigt werden sollten, und das Aufzeigen von Spannungsfeldern zwischen unterschiedlichen Interessenvertretern und deren Anforderungen, welche die Ausprägungen dieser Dimensionen beeinflussen.
- (ii) Eine Anwendung der aufgezeigten Dimensionen im Rahmen der Erweiterung aktueller bFIM-Systeme zu UCAID demonstriert das Potenzial eines systematischen Entwurfs.
- (iii) Eine Diskussion verschiedener Ansätze zur Initiierung des Informationsaustauschs

auf der Basis eines Prototypen von UCAID.

2 Dimensionen

Im Folgenden werden unterschiedliche Dimensionen aufgezeigt, welche es bei der Verwaltung identitätsbezogener Informationen in verteilten Systemen zu berücksichtigen gilt. Bei den unterschiedlichen Ausprägungen dieser Dimensionen handelt es sich um konzeptionelle Möglichkeiten, die unabhängig davon zu betrachten sind, ob sie in bestehenden FIM-Systemen bereits umgesetzt wurden. Für die einzelnen Dimensionen wird dargelegt, welche Anforderungen an ein System die jeweilige Dimension beeinflussen kann und in welchem Spannungsfeld diese Anforderungen und damit auch die verschiedenen Interessenvertreter – Benutzer, IdP und SP – stehen. Die Beschreibung der einzelnen Dimensionen soll hierbei keinesfalls Anspruch auf Vollständigkeit erheben, sondern vielmehr als möglicher Ausgangspunkt einer Taxonomie gesehen werden, mit deren Hilfe eine systematische Einordnung und Bewertung von Ansätzen zur Verwaltung von Identitätsinformationen möglich wäre. Folgende Dimensionen können unterschieden werden:

Speicherort – Wer hält die zu verwendenden Daten vor? - Eine Alternative besteht darin, dass *der Benutzer* die Daten selbst verwaltet. In diesem Fall hat der Benutzer die volle Kontrolle über seine Identitätsinformationen und kann diese, sobald sich eine Änderung ergibt, unmittelbar aktualisieren. Hieraus folgt, dass in die Validität dieser Informationen seitens eines SP, der die Informationen später erhält, nur ein geringes Maß an Vertrauen gesetzt werden kann.

Hält *der IdP* hingegen die Informationen vor, so kann dieser die vom Benutzer bereitgestellten Informationen überprüfen und für die Validität der Informationen zu einem gewissen Grad bürgen. Dies ermöglicht es dem SP, in die Validität der vom IdP bereitgestellten identitätsbezogenen Informationen ein bestimmtes Maß an Vertrauen setzen zu können, was insbesondere dann wichtig ist, wenn basierend auf diesen Informationen der Zugriff auf den SP kontrolliert werden soll.

Der SP kann die identitätsbezogenen Informationen des Benutzers zusätzlich selbst vorhalten. Falls die Abfrage der Informationen von einem IdP zum Zeitpunkt der Dienstbringung nicht möglich ist, muss der SP die Informationen lokal speichern. Für eine lokale Speicherung können noch weitere Gründe existieren, z. B. die Verfügbarkeit der Informationen bei Ausfall von Netzverbindungen. Bei einer lokalen Speicherung muss jedoch gewährleistet werden, dass die lokal vorgehaltenen Informationen bei einer Änderung aktualisiert werden.

Informationsaustausch – Wer initiiert den Informationsaustausch? - Der Informationsaustausch kann auf Initiative *des Benutzers* erfolgen, indem er die Informationen direkt an den SP weitergibt. Dies stellt jedoch sehr hohe Anforderungen an den Benutzer, denn es setzt voraus, dass der Benutzer einen Überblick über sämtliche Replikate seiner identitätsbezogenen Informationen hat. In Anbetracht der steigenden Anzahl an genutzten Diensten gestaltet sich dies zusehends schwierig.

Eine weitere Möglichkeit ist ein Push-basierter Ansatz, der vorsieht, dass *der IdP* dem

SP Informationen zusendet. Während dieser Ansatz eine schnelle Aktualisierung der Informationen nach einer Änderung ermöglicht, setzt er voraus, dass der IdP sowohl die Verwaltung der Empfänger als auch Maßnahmen für die Zugriffskontrolle realisiert.

Die dritte Alternative bildet ein Pull-basierter Ansatz, bei dem *der SP* identitätsbezogene Informationen anfordert. Dies kann bei Bedarf oder in periodischen Abständen geschehen. Das Intervall zwischen zwei Anfragen des SP wird hierbei u. a. durch die Konsistenzanforderungen des SP selbst bestimmt und wirkt sich wiederum auf die im Netzwerk generierte Last aus.

Auch hybride Ansätze wären möglich: Der IdP könnte bspw. den SP unmittelbar über die Änderung von Informationen benachrichtigen, ohne hierbei die geänderten Informationen mitzusenden. Der SP könnte aktuelle Informationen dann bei Bedarf anfordern.

Discovery – Wie wird festgelegt, woher die Daten stammen? - Falls die identitätsbezogenen Informationen eines Benutzers von einem oder mehreren IdPs vorgehalten werden, muss für den konkreten Bezug der Informationen festgelegt werden, von welchem IdP die Informationen geholt werden sollen.

Dies kann zum einen *durch den Benutzer* geschehen, der den zu verwendenden IdP gegenüber dem SP angibt; dieser Ansatz stellt die Bestimmung des Benutzers über den Speicherort seiner Informationen in den Vordergrund. Gerade wenn der SP die identitätsbezogenen Informationen des Benutzers von verschiedenen IdPs beziehen muss, kann die Angabe aller relevanten IdPs in einem hohen Aufwand für den Benutzer resultieren. Es könnte daher im Sinne der Benutzbarkeit von Vorteil sein, die Lokalisierung der Informationen einem Drittanbieter zu überlassen, der als Verzeichnisdienst fungiert. Der Benutzer würde bei diesem Verzeichnisdienst angeben, von welchen IdPs seine Informationen bezogen werden können. Der SP könnte diese Information beim Verzeichnisdienst erfragen und im Anschluss die entsprechenden IdPs kontaktieren.

Ein eher dienstorientierter Ansatz sieht vor, dass der zu verwendende IdP *durch den SP* von vornherein festgelegt oder dem Benutzer eine eingeschränkte Menge an IdPs zur Auswahl gestellt wird. Insbesondere wird dies notwendig, wenn der SP validierte Informationen benötigt und deshalb Informationen von einem IdP nur dann akzeptiert, wenn zu diesem eine Vertrauensstellung besteht. In der Frage nach der Festlegung der Herkunft identitätsbezogener Informationen wird das Spannungsfeld ersichtlich, in dem sich Benutzerkontrolle und Validierung von Informationen befinden.

Autorisation – Wann findet die Autorisation statt? - Ein Grundgedanke des Datenschutzes besteht darin, dem Benutzer die Kontrolle über die Weitergabe seiner identitätsbezogenen Informationen zu ermöglichen. Zu diesem Zweck muss der Benutzer u. a. entscheiden können, welche Informationen an welchen SP weitergegeben werden sollen, er muss also den SP für den Zugang zu den Informationen autorisieren.

Dies kann einerseits *unmittelbar vor jedem Zugriffsversuch* des SP auf die Informationen geschehen, was eine Interaktion mit dem Benutzer unmittelbar vor dem Zugriffsversuch erforderlich macht. Dies erlaubt einerseits eine feingranulare Autorisierung durch den Benutzer, kann allerdings angesichts der wiederholt erforderlichen Anwendung desselben Schrittes die Benutzbarkeit des Systems beeinträchtigen.

Alternativ dazu kann der Benutzer die Autorisation bereits *im Vorfeld des Zugriffs* vornehmen. Zu diesem Zweck kann der Benutzer bspw. Richtlinien definieren, die dann als Basis einer späteren Zugriffskontrollentscheidung herangezogen werden.

Zugriffskontrolle – Wo bzw. durch wen findet Zugriffskontrolle statt? - Die Kontrolle des Zugriffs kann *durch den Client des Benutzers* selbst vorgenommen werden; in diesem Fall muss jede Übermittlung identitätsbezogener Informationen unter Einbeziehung des Clients des Benutzers erfolgen. Dies ermöglicht ein hohes Maß an Kontrolle des Benutzers über seine Informationen und die Protokollierung der Verteilung der Informationen. Auf der anderen Seite muss bei diesem Ansatz der Client des Benutzers zum Zeitpunkt der Weitergabe über das Netzwerk erreichbar sein, was eine automatisierte Aktualisierung der identitätsbezogenen Informationen erschwert.

Die Kontrolle des Zugriffs auf die identitätsbezogenen Informationen des Benutzers könnte alternativ *durch den IdP* erfolgen. Bei diesem Ansatz müsste der Client des Benutzer für die Weitergabe identitätsbezogener Informationen nicht erreichbar sein. Allerdings wären in diesem Fall die Kontrollmöglichkeiten des Benutzers über die Weitergabe seiner Informationen eingeschränkt, woraus ein höheres Maß an Vertrauen in den IdP resultiert.

Eine weitere Möglichkeit würde darin bestehen, die Zugriffskontrolle einem *durch einen Drittanbieter bereitgestellten Dienst* zu überlassen, der als Zwischenstation in jede Übermittlung von Informationen involviert ist. Seitens des Benutzers müsste in diesen Dienst ebenfalls ein gewisses Maß an Vertrauen gesetzt werden. Es gilt zu beachten, dass der Dienst, der den Zugriff kontrolliert, nicht zwingend auch die Zugriffsentscheidung treffen muss. Diese Aufgabe kann ein weiterer Dienst wahrnehmen, der vom kontrollierenden Dienst vor jedem Zugriff kontaktiert wird.

3 Illustrative Anwendung der Dimensionen

Im Folgenden wird aufgezeigt, wie die in Abschnitt 2 beschriebenen Dimensionen zur systematischen Weiterentwicklung eines Systems oder Ansatzes zur Verwaltung identitätsbezogener Informationen genutzt werden können. Zu diesem Zweck wird bFIM als Ansatz vorgestellt. Danach wird zunächst abstrakt und anschließend exemplarisch veranschaulicht, wie mit Hilfe der Dimensionen ausgehend von einer neuen Anforderung an das System der Entwurf einer entsprechenden Erweiterung systematisiert werden kann.

Das Ziel, die Kontrolle des Benutzers über seine identitätsbezogenen Informationen zu gewährleisten, wird insbesondere von bFIM-Systemen angestrebt [MR08]. In bFIM-Systemen legt der Benutzer diese Informationen initial bei einem IdP ab; hierbei wird u. U. die Validität der Informationen überprüft. Der Zugriff auf die vom IdP vorgehaltenen Informationen ist in bFIM-Systemen ausschließlich dem Benutzer erlaubt. Um die Informationen an einen SP weiterzugeben, muss der Benutzer sich gegen den IdP authentisieren, die Informationen anfordern und an den SP weiterleiten. Der Client des Benutzers ist also stets in die Weitergabe der Informationen involviert. Dieser von aktuellen bFIM-Systemen verfolgte Ansatz soll im Folgenden exemplarisch unter Anwendung der Dimensionen erweitert werden.

Der Nutzen der vorgeschlagenen Dimensionen ergibt sich daraus, dass die Auswirkungen der unterschiedlichen Ausprägungen der Dimensionen auf die Eigenschaften von Ansätzen zur Verwaltung identitätsbezogener Informationen in Abschnitt 2 qualitativ beschrieben sind. Umgekehrt können also ausgehend von einer Anforderung an einen Ansatz die Dimensionen abgeleitet werden, deren Änderung zum Umsetzen der Anforderung beiträgt. Somit können auf systematische Art und Weise sowohl neue Ansätze entworfen als auch – wie im Folgenden veranschaulicht wird – bestehende Ansätze erweitert werden.

Zunächst wird hierzu der grundlegende Ansatz von bFIM-Systemen bezüglich der Dimensionen eingeordnet. Der *Speicherort* der identitätsbezogenen Informationen kann in bFIM-Systemen variieren: Die identitätsbezogenen Informationen können durchaus vom Benutzer selbst vorgehalten werden, zusätzlich ist es jedoch üblich, dass die Informationen bei einem oder mehreren IdPs abgelegt werden. Es steht dem SP frei, die Informationen zusätzlich lokal vorzuhalten. Der *Informationsaustausch* kann in bFIM-Systemen ausschließlich vom Benutzer initiiert werden. In bFIM-Systemen obliegt der Vorgang des *Discovery* letztendlich dem Benutzer. Dieser entscheidet selbst, von welchem IdP die identitätsbezogenen Informationen bezogen und dem SP übergeben werden sollen. Die Auswahl der möglichen IdPs kann jedoch durch den SP beschränkt werden, insbesondere wenn der SP validierte Informationen benötigt und aus diesem Grund eine Vertrauensstellung zum IdP unterhalten muss. Die *Autorisation* findet in bFIM-Systemen unmittelbar vor dem Dienstzugriff statt, indem der Benutzer die Authentisierung gegen den IdP vornimmt oder ablehnt. Dies verhindert gleichzeitig, dass identitätsbezogene Informationen zum SP übertragen werden, somit nimmt also der Benutzer bzw. der Client des Benutzers die *Zugriffskontrolle* vor.

Eine Unzulänglichkeit von bFIM-Systemen besteht darin, dass eine Weitergabe von identitätsbezogenen Informationen ohne unmittelbare Beteiligung des Benutzers nicht möglich ist. Dies bedeutet, dass ein SP für die Aktualisierung lokal vorgehaltener Informationen auf die Initiative des Benutzers angewiesen ist. Die Aktualität dieser Informationen ist somit nicht zu jeder Zeit gewährleistet, wie in [HKDH11] gezeigt wurde. Dies erweist sich dann als problematisch, wenn die Dienstleistung durch den SP unabhängig von der letzten Aktualisierung der Informationen durch den Benutzer erfolgen muss und hierfür aktuelle Informationen benötigt werden, um bspw. eine Zugriffskontrolle vorzunehmen.

Dieses Problem kann adressiert werden, indem an bFIM-Systeme die zusätzliche Anforderung gestellt wird, dass eine Übermittlung identitätsbezogener Informationen auch ohne unmittelbares Zutun des Benutzers und ohne Involvierung des Client des Benutzers möglich sein muss. Basierend auf dieser Anforderung kann aus den in Abschnitt 2 dargelegten Dimensionen Nutzen gezogen werden: Da die Anforderung einer Übermittlung ohne Zutun des Benutzers von den Dimensionen *Informationsaustausch*, *Autorisation* und *Zugriffskontrolle* abhängig ist, muss folglich der Entwurf einer entsprechenden Erweiterung von bFIM-Systemen diese Dimensionen in den Fokus stellen.

Diese systematische, an den Dimensionen orientierte Herangehensweise erleichterte den Entwurf von UCAID, welches bereits in [HKDH10] und [HKDH11] umfassend dargelegt wurde. Im folgenden Abschnitt wird UCAID kurz vorgestellt.

3.1 Überblick über UCAID

UCAID erweitert bFIM-Systeme um die Möglichkeit einer Weitergabe identitätsbezogener Informationen, ohne dass der Client des Benutzers involviert ist, während gleichzeitig die Kontrolle des Benutzers über die Informationen gewahrt bleibt. Der vorgeschlagene Ansatz fügt dem System einen zusätzlichen Dienst mit der Bezeichnung „Identity Delegate“ hinzu, welcher stellvertretend für den Benutzer den Bezug identitätsbezogener Informationen durch SPs überwacht, insbesondere dann, wenn der Benutzer offline ist. Hierbei stellt der Identity Delegate zwei Hauptaspekte der Benutzerzentrierung sicher: Zum einen vertritt der Identity Delegate den Benutzer im Rahmen der Kontrolle der Weitergabe der Informationen, indem er benutzerdefinierte Richtlinien auswertet, sobald ein SP Informationen abrufen möchte. Zum anderen tritt der Identity Delegate hinsichtlich des Datenflusses an die Stelle des Client des Benutzers, d. h., die Informationen passieren den Identity Delegate anstatt den Client des Benutzers. UCAID verlangt für die Kontrolle der Weitergabe der Informationen kein manuelles Zutun des Benutzers zum Zeitpunkt der Anfrage. Stattdessen übt der Benutzer die Kontrolle über die Weitergabe aus, indem er Richtlinien definiert, die die Grundlage für eine *Automatisierung* der Kontrolle darstellen.

Der Benutzer vertraut dem Identity Delegate dahingehend, dass dieser identitätsbezogene Informationen entsprechend der benutzerdefinierten Richtlinien bezieht und weitergibt. Nach dem Prinzip des „Least Privilege“ werden dem Delegate jedoch nur eingeschränkte Rechte gewährt. Hierdurch wird ein unbeschränkter Zugang zu den Informationen des Benutzers verhindert und damit das erforderliche Maß an Vertrauen seitens des Benutzers in den Delegate reduziert. Insbesondere setzt der Ansatz voraus, dass die IdPs so konfiguriert werden, dass die identitätsbezogenen Informationen vor der Herausgabe signiert und für den anfragenden SP verschlüsselt werden. Das hierfür notwendige Maß an Vertrauen von SP und Benutzer in den IdP wird im föderativen Identitätsmanagement generell vorausgesetzt. Auf diese Weise agiert der Identity Delegate als „Identity Relay“ [CI09], welches identitätsbezogene Informationen lediglich abholt und weiterleitet, aber nicht selbst lokal ablegt. Folglich ist ein Angreifer, der Kontrolle über den Identity Delegate erlangt, nicht in der Lage, identitätsbezogene Informationen des Benutzers zu lesen, zu speichern oder zu verändern.

UCAID modifiziert bFIM-Systeme insbesondere hinsichtlich der Dimensionen *Informationsaustausch*, *Autorisation* und *Zugriffskontrolle*, wie anhand der in UCAID erforderlichen Schritte veranschaulicht werden kann: Im Vorbereitungsschritt muss der Delegate, der als von einem Drittanbieter bereitgestellter Dienst entworfen wurde, für eine Nutzung vorbereitet werden. Im nächsten Schritt, dem Registrierungsschritt, macht der Benutzer den Delegate mit seinen Identitäten bei einem oder mehreren IdPs bekannt. Im Autorisierungsschritt ermächtigt der Benutzer einen SP dazu, über den Identity Delegate identitätsbezogene Informationen von den IdPs zu beziehen, indem der Benutzer SP-spezifische Zugriffsrichtlinien definiert. Der Autorisierungsschritt kann wiederholt werden, um weitere SPs für den Bezug der Informationen zu autorisieren. Diese Ausprägung der *Autorisation* unterscheidet sich von der Ausprägung aktueller bFIM-Systeme, die für jede Weitergabe der Informationen eine erneute Autorisation verlangen. Der letzte und wesentliche Schritt ist der Aktualisierungsschritt, in dem der SP vom Identity Delegate Informationen an-

fordern und ggf. beziehen kann. Die Initiative zum *Informationsaustausch* geht also vom SP aus und nicht mehr vom Benutzer. Die für den Schutz der identitätsbezogenen Informationen notwendige *Zugriffskontrolle* wird vom Identity Delegate vorgenommen. Somit erfordert dieser letzte Schritt kein manuelles Eingreifen durch den Benutzer.

3.2 Bewertung hinsichtlich Benutzerzentrierung und Föderation

Im Vergleich zu aktuellen Ansätzen des bFIM wurde die Ausprägung in den Dimensionen *Informationsaustausch*, *Autorisation* und *Zugriffskontrolle* angepasst. Hinsichtlich der Dimension *Informationsaustausch* wurden durch den Ansatz die Anforderungen des SP adressiert, da zum Erreichen einer verbesserten Aktualität der Identitätsinformationen eines Benutzers beim SP der Zeitpunkt zur Abfrage aktueller Informationen durch den SP bestimmt werden muss. Des Weiteren wurde die Ausprägung der Dimension *Autorisation* adaptiert. Diese Anpassung kann jedoch nicht als Schwächung der *Benutzerzentrierung* angesehen werden, da sich argumentieren lässt, dass durch UCAID der Benutzer nicht ständig nach seiner Zustimmung gefragt wird, sondern der Benutzer durch das Einrichten der Richtlinien dauerhaft festlegen kann, welcher Dienst welche Berechtigungen besitzt. Diese Ausprägung kann zu einer Reduktion der Fehleranfälligkeit führen, da der Benutzer die Einstellung der Richtlinien weniger häufig durchführt als die Erteilung der Berechtigung in aktuellen bFIM-Systemen, und somit möglicherweise eine größere Sorgfalt auf die Einstellung der Richtlinien verwendet. Zusätzlich zur *Autorisation* wird durch UCAID auch die *Zugriffskontrolle* nicht mehr durch den Benutzer, sondern durch den Identity Delegate vorgenommen. Um das Risiko einer unberechtigten Weitergabe von Identitätsinformationen zu reduzieren und somit auch das notwendige Vertrauen zu verringern, welches der Benutzer in den Identity Delegate setzen muss, wird in UCAID auf Delegationsmechanismen zurückgegriffen. Hierdurch wird es möglich, dass der Identity Delegate streng nach dem Prinzip des „Least Privilege“ nur die Rechte erhält, die zum Ausführen seiner Aufgabe notwendig sind.

Die beschriebenen Dimensionen können wiederum angewandt werden, um Schwachstellen in UCAID zu untersuchen. Exemplarisch sollen in diesem Abschnitt Ansätze zur Verringerung der von UCAID erzeugten Netzwerklast verglichen und bewertet werden. Wie in Abschnitt 2 skizziert, wird die generierte Netzwerklast wesentlich von der Ausprägung des *Informationsaustauschs* bestimmt; somit kann diese Dimension den Ausgangspunkt für eine mögliche Modifikation von UCAID bilden.

In UCAID wird der Informationsaustausch durch den SP initiiert, d. h., UCAID beruht auf einem Pull-Ansatz. Da die Änderungsrate eines Großteils identitätsbezogener Informationen eher gering ist, können in einem Pull-basierten Ansatz potenziell Daten angefragt werden, selbst wenn sich keinerlei Änderungen ergeben haben. Folglich ist in diesem Szenario bei einem Pull-Ansatz die Netzwerklast tendenziell höher als in einem Push-Ansatz. Eine mögliche Umstellung von UCAID auf einen Push-Ansatz wird im Folgenden beschrieben und mit dem bestehenden Pull-Ansatz verglichen. Die vorgestellten Ansätze werden des Weiteren gemessen, um deren Einfluss auf die Netzwerklast zu ermitteln. Da diese Messungen nur exemplarisch für eine bestimmte Konfiguration vorgenommen werden können,

wird im Folgenden angenommen, dass der Identity Delegate 3 Anfragen pro Sekunde beantworten muss. Die Anfragerate von 3 Anfragen pro Sekunde wurde basierend auf Kennwerten festgelegt, die einem produktiven IDM-System entnommen sind. Hieraus resultiert eine Last von ca. 1700 kbit/s (vgl. Abb. 1).

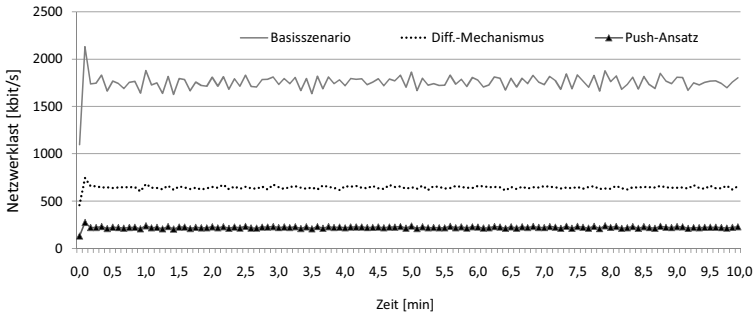


Abbildung 1: Netzwerklast in einem exemplarischen Szenario gemessen zwischen dem Identity Delegate und den anfragenden SPs.

Da eine wesentliches Ziel bei der Umsetzung von UCAID war, dass IdPs möglichst wenig angepasst werden müssen, wurde in den im Folgenden vorgeschlagenen Ansätzen zur Reduktion der Netzwerklast von einer Reduktion der Last zwischen dem Identity Delegate und den IdPs abgesehen. Insbesondere wird nicht in Betracht gezogen, dass ein IdP Änderungen an den Identity Delegate aktiv mitteilt. Es wird somit versucht, die Kommunikation zwischen dem Identity Delegate und den Diensteanbietern effizienter zu gestalten. Da die Netzwerklast zwischen dem Identity Delegate und den IdPs nicht reduziert wurde, werden alle Messungen des Kommunikationsaufkommens auf die Kommunikation zwischen dem Identity Delegate und den anfragenden SPs beschränkt. Folgende Ansätze zur Reduktion der Netzwerklast wären denkbar:

“Diff-Mechanismus” – ein Diff-Mechanismus erlaubt zu ermitteln, ob sich seit der letzten Anfrage eine Änderung ergeben hat oder nicht. Der Identity Delegate könnte somit bei einer Anfrage bei dem IdP den Zeitpunkt der letzten Änderung nachfragen. Nur wenn sich an den Informationen seit der letzten Anfrage des SP eine Änderung ergeben hat, würde er die geänderten Informationen kommunizieren. Andernfalls würde er lediglich übermitteln, dass sich seit der letzten Anfrage keine Änderung ergab, wodurch sich eine geringer Auslastung des Netzwerks ergibt. Die konkrete Netzwerklast ist direkt mit der Änderungsrate der identitätsbezogenen Informationen eines Benutzers korreliert. Bei einer Änderungsrate von bspw. einer Änderung pro Woche, würde sich durch den Einsatz eines Diff-Mechanismus die Netzwerklast auf ca. 600 kbit/s reduzieren lassen (vgl. Abb. 1).

Push-Ansatz des Identity Delegate – ein Erweiterung hinsichtlich des Einsatzes eines Diff-Mechanismus ist es, dass SPs nicht weiter den Identity Delegate regelmäßig abfragen, sondern der Identity Delegate die SPs über Änderungen informiert. In diesem Fall kann bei einer exemplarischen Änderungsrate von einer Änderung pro Woche die Netzwerklast auf ungefähr 200 kbit/s gesenkt werden.

4 Zusammenfassung und Ausblick

Um in Anbetracht einer von Einzelaspekten geprägten Sicht auf Ansätze zur Verwaltung identitätsbezogener Informationen die Vergleichbarkeit zu ermöglichen, wurden in der vorliegenden Arbeit Dimensionen aufgezeigt, auf deren Basis ein systematischer Entwurf und Vergleich dieser Ansätze möglich ist. Obwohl diese Dimensionen als erweiterbar anzusehen sind, konnte eine illustrative Anwendung dieser Dimensionen, in der der aktuelle bFIM-Ansatz zu UCAID erweitert wurde, die Möglichkeit zur systematischen Entwicklung aufzeigen.

Wie in dieser Arbeit dargelegt wurde, bewegen sich Ansätze zur Verwaltung von Identitätsinformationen in Spannungsfeldern, die bei der Beantwortung der eingangs gestellten Frage „Wie kann ein Benutzer die Kontrolle über seine Identitätsdaten bewahren?“ stets bedacht werden müssen. Insbesondere bedeutet dies, dass neben der Kontrolle des Benutzers über seine Informationen und der Benutzerfreundlichkeit auch die Anforderungen der Dienstanbieter, bspw. hinsichtlich der Aktualität der Informationen, adäquat erfüllt werden müssen.

Das Ziel zukünftiger Arbeiten wird es sein, die in dieser Arbeit vorgestellten Dimensionen weiter zu spezifizieren und zu verfeinern, um letztendlich eine Taxonomie von Ansätzen zur verteilten Verwaltung von Identitätsinformationen aufzustellen, welche eine Einordnung und einen Vergleich dieser Ansätze auf systematische Art und Weise erlaubt.

Literatur

- [CI09] D. W. Chadwick and G. Inman. Attribute Aggregation in Federated Identity Management. *IEEE Computer*, 42(5):33–40, 2009.
- [HKDH10] T. Hoellrigl, H. Kühner, J. Dinger, and H. Hartenstein. User-Controlled Automated Identity Delegation (Short Paper). In *6th IEEE/IFIP International Conference on Network and Service Management (CNSM)*, pages 230–233, 2010.
- [HKDH11] T. Hoellrigl, H. Kühner, J. Dinger, and H. Hartenstein. Extension for Information Card Systems to Achieve User-Controlled Automated Identity Delegation. In *1st IEEE/IFIP Workshop on Managing Federations and Cooperative Management (Man-Fed.CoM)*, 2011.
- [MMCvM10] M. Machulak, E. Maler, D. Catalano, and A. van Moorsel. User-Managed Access to Web Resources. Technical report, University of Newcastle upon Tyne: Computing Science, 2010.
- [MR08] Eve Maler and Drummond Reed. The Venn of Identity: Options and Issues in Federated Identity Management. *IEEE Security and Privacy*, 6(2):16–23, 2008.
- [OAu10] OAuth Core 1.0, 2010. <http://oauth.net/core/1.0a>.

How I and others can link my various social network profiles as a basis to reveal my virtual appearance

Sebastian Labitzke, Jochen Dinger, Hannes Hartenstein

sebastian.labitzke@kit.edu, jochen.dinger@fiducia.de, hannes.hartenstein@kit.edu

Abstract: Today, many Internet users take part in online social networks (OSNs) like Facebook, Xing, LinkedIn, MySpace, studiVZ, and others. However, the motives for participation are multifaceted and stand in contrast to the motivation not to join OSNs because of the potential danger with respect to privacy aspects. This risk is already perceived by some OSN users but still underestimated. The naive handling with personal data inside such networks combined with many possibilities to link several OSN profiles of a single user paves the way for third-parties to gather the comprehensive virtual appearance of a person. Therefore, to turn the tables, such correlations of information might be useful to demonstrate to a user a measure of how linkable his personal data is, given the current exposure of his public data in OSN profiles. In this paper we point out some astonishing simple ways to automatically link data from different OSNs. As we want to and have to act compliant to the German data protection laws, we analyze corresponding requirements and sketch our concept for compliant statistical sampling. We further show preliminary results regarding successful profile correlations based on, even under legal and technical constraints, extracted friends lists.

1 Introduction

The online social network (OSN) Facebook is actively used by over 500 million people (logged in within the last 30 days) and grew by 500 percent within the last two years¹. The number of members of many other social networks is growing as well (e.g. [MKG⁺08]), users are typically member of more than a single OSN, and in every OSN people publish a lot of personal data. For instance, Facebook claims that an average user uploads and shares 90 pieces of content each month. The main reasons for joining OSNs and sharing such amount of data are the possibility to meet so called “friends” (and strangers as well), the chance to keep in touch with those people, and not least the opportunity to express oneself through uploaded photos, videos, comments, or other small contributions. Also the asynchronous means of communications make OSNs attractive for many people.

Quite clearly, thoughtless sharing of personal data bears certain risks. The obvious question is: why do many users not think about these risks while providing personal data? Answers can be found in behavioral psychology with respect to rewards and punishments of human actions. The exposure of personal data ensures direct rewards in terms of enabling the user to participate and take part in the OSN. The more data, and consequently content, a user shares, the better he can be located and identified inside the network and

¹<http://www.facebook.com/press/info.php?timeline>

his profile will arise more interest. Reactions in terms of comments on the user's new content intensify these rewards. Furthermore, by sharing content about oneself an individual distinctive urge for self-representation may be satisfied, which can also be a desired reward for users. Thus, the motivation to publish content is increased by the expectation of benefits. Psychologists refer to this as "positive reinforcement" of the motivation to do something [Zim04].

In order to develop users' risk awareness regarding privacy aspects, a negative reinforcement, i.e., an increased motivation not to share personal data in the sense of conditioning of users, is necessary. Referring to this, it is known that the time gap between an action and a consecutive set of rewards and punishments, as well as their nature and severity, shapes future behavior. This time gap is known as "deferred gratification" ([MM83], [Mis74]). In the context of OSNs, rewards occur on a short and punishments, if any at all, on an often longer time scale. As a consequence, punishments and therefore sustainable recognitions of the potential of danger are less effective on the motivation to share personal data in OSNs than the above discussed rewards a user will get promptly. Such recognitions or rather a conditioning can be achieved by different approaches. In general, promising options are the intensification of the consequences of an action and/or the reduction of the time gap between the action (here: providing personal data) and the occurrence of punishments. Psychologists refer to the conditioning by regulated rewards and punishments as "operant conditioning". However, it is not clear which punishments should be provoked to achieve an effect on users' behavior in OSNs. Furthermore, punishments would unnecessarily cause detriments to users as well as OSNs would not endorse any sanctions on publishing data as this stands obviously in contrast to their business model.

Meanwhile, discussions and reports about privacy with respect to OSNs in the media are urging more users to make use of privacy settings that are provided by OSNs to hide parts of a profile from strangers. This is also confirmed by a German survey, in which 1,208 teenagers between 12 and 19 years were asked about their behavior in OSNs [jim10]. In the recent survey two thirds confirmed that they hide parts of their profiles from strangers, while in an earlier study (2009) less than a half of those interviewed stated that they make use of privacy settings. In 2008 only 25 percent of Facebook users made use of restricted privacy settings [KW08]. While more and more OSN members have adjusted their privacy settings, the fact that a large number of them are still unduly generous in sharing personal data applies to most OSNs. In particular, many users still do not hide parts of their OSN profiles or do not hide every published information from strangers (cf. Section 2). Moreover, people are often member of more than a single OSN, even though Facebook acquires a remarkable market share in recent time. Users' multiplicity of profiles are comprehensible due to the different objectives of current OSNs (business and private networks, networks for fellow students etc.). Furthermore, Torkjazi et al. show that OSNs usually get into a downturn of interest after a major increase of members [TRW09] which prompt users to register in additional OSNs. According to that, if a person has more than one OSN profile, it is often easy to merge these profiles on the basis of provided information. Such linking of profiles can be used to gather and associate various shared information [MV09] by third-parties, but also to illustrate to a user how his comprehensive virtual appearance looks like despite possibly adjusted privacy settings. In contrast to inapplicable operant

conditioning, such an illustration might be a far better conditioning approach to encourage users to hide their personal data in OSNs. The aim is to demonstrate the potential of danger due to the disclosure of measures of how easy information of OSN profiles can be associated.

Furthermore, such possibility to reveal users' own virtual appearance by linking OSN profiles might constitute a basis for a novel identity management system with which users will be able to monitor the proliferation of their published information on their own in the future. On the contrary, in enterprise identity management systems the "flow" and processing of identity related information is typically monitored and governed by various overview boards. While the publishing of personal information in OSNs is essentially at the users discretion, the flow of such information and their appearance in several systems can not be monitored or further managed by users, today.

In this paper, we take a first step in the direction of such an identity (self-)management component. Therefore, we assess which kind of personal information is published by users and discuss possibilities to correlate users' profiles from different OSNs based on this data. More precisely, we unfold the availability of personal attributes in different OSNs that can be used to link OSN profiles. Our investigations demonstrate that many profiles still provide sufficient information to establish such links between OSN profiles. This is particularly critical if a user makes use of restrictive privacy settings in one OSN and non-restrictive settings in another. While the restrictive settings shall help to protect provided data, the settings might be ineffective due to the linkability of both profiles. The still existing risk of an attack on privacy can apparently be overlooked by users. Additionally, we show how crucial friends lists are in order to link OSN profiles. Our statistical analysis on the linkability of profiles on the basis of friends list reveals the "fingerprint" character of such individual lists of names. To get such specific knowledge we performed statistical analysis on the users' data without archiving any reference to OSN profiles or sufficient information to identify a natural person. According to our analysis of German Data Protection Act and to the best of our knowledge, the study (presented in Section 5) is compliant to the German law and ensures users' privacy. Hence, in this paper a "lower bound" is made explicit that indicates what is possible by correlation of OSN profiles without breaking the law. However, it is easy to imagine the potentialities for people who do not care about legal aspects. Such people are merely faced with some technical measures implemented by OSNs to avoid software-based extractions of personal data by third-parties. In the course of this paper we additionally point out the deficiency of such measures that in turn should not prevent automatic data extraction.

Section 2 presents an overview of related work that addresses the amount of voluntarily published data in OSNs and the linking of profiles. Section 3 discusses restrictions for studies on personal data of OSN profiles on the basis of the German Data Protection Act and introduces requirements to act compliant. Section 4 points out technical measures of OSNs to prevent extractions of personal data and assesses their effectiveness. We then look at the concept for an implementation to carry out the study, at statistical results concerning the concrete data users provide in OSNs, and at a preliminary estimated quality of profile correlations on the basis of friends lists in Section 5. Section 6 concludes the paper.

2 Related work

In the past, several researchers investigated the publicly available amount of information inside OSNs. As early as 2005 the use of privacy settings of 4,000 Carnegie Mellon University students were determined by Gross and Acquisti [GA05]. They discovered that 89% of analyzed users provided their full real names and 90% a profile image. Furthermore, 80% of analyzed profiles contained information that can be used to identify a person. Merely 1.2% of users configured that they cannot be found by using the OSN search and only 0.06% restricted the visibility of their profile. In 2007 30,000 Facebook profiles were analyzed by Lampe et al. [LES07]. At that time, already 19% made use of privacy settings but in average 59% of provided fields inside a profile were filled with personal information. These studies combined with parts of our results presented in Section 5 show the evolution of risk awareness and sensitivity of OSN members.

In [KW08] and [KW10a] Krishnamurthy and Wills analyzed the default configurations of privacy settings, which information can be hidden from strangers, and how many users adjusted the privacy settings in 12 OSNs. They pointed out that in most OSNs personal data is classified in categories. Users are able to adjust the visibility of these categories instead of more fine grained access configurations. Furthermore, in [KW08] it is shown that only 1% of Twitter users, 25% of Facebook users, and 21% of MySpace users have adjusted their privacy settings whereas in [KW10a] the authors disclose the technical possibilities for third-parties to gather personal information of OSN users. In contrast to our investigation, they analyzed a way to get information about OSN profiles of a visitor of a website by tracking him surfing through different websites and OSNs. For a follow-up study and on the basis of the knowledge of their previous studies, the same authors compared 13 mobile OSNs in 2010 [KW10b]. They demonstrated that friends lists are available in six mobile variants of OSN websites, in other six by default and in one network friends lists are never publicly available. These three papers show that a lot of information (especially friends lists) is still accessible for third-parties. This in turn underpins our thesis that average OSN profiles still provide enough information to link several profiles of a single person.

Besides users behavior in OSNs with respect to the amount of publicly available personal data, the linkability of OSN profiles has been examined. The authors of [MV09] demonstrated which criteria are suitable to effectively identify a user in an OSN. Therefore, they extracted several attributes of OSN profiles and, afterwards, they used these attributes to search a profile inside another network. With this study the authors pointed out that 25.2% of all analyzed Facebook profiles overlapped with a single MySpace profile. In contrast to our study, the authors used as much information as needed to identify a person inside a network, whereas we quantify the potential of single attributes regarding their suitability to link one OSN profile to another profile in a different OSN. Thereby, our aim is to identify information in OSN profiles that is more worthy of protection than others. The aim of the authors work was to provide a system to locate individuals in Facebook and MySpace as well as to figure out how many users are member of both of these OSNs. Furthermore, the authors stated that friends lists of two profiles of the same natural person just rarely overlap. We show in this paper that this more or less rare overlap is yet sufficiently large to distinguish it from overlaps of profiles from different people.

3 Compliance requirements on statistical analysis

For a tool that should enhance users' privacy as well as for an underlying study, compliance with privacy regulation is a must. Therefore, this section discusses legal requirements which have to be met for statistical sampling that ensures privacy of OSN members.

Derived from German constitution (Art.2x1 in conjunction with Art.1x1, GG) and German Data Protection Act (BDSG) the right to privacy and the right to self-determination with respect to personally identifiable information are established. According to the BDSG §4.3 and §4a.1 personal information may only be used for a desired purpose and only an owner of information, respectively the OSN user himself, can approve data processing by third-parties. Users' data inside OSNs are publicly available to take part of the network and activities inside the network. Such data is not allowed to be used by third-parties for further processing. Some OSNs explicitly state in their terms and conditions that automated extractions of data are prohibited to protect OSN members.

According to our analysis of BDSG, it is not possible to legitimize any extraction, processing, and archiving of users' profiles and contained personal data without a previously stated agreement of any affected user (permission facts). Neither §40 (BDSG) is applicable, which regularizes the processing and use of personal data by research institutions, because it is required that the research institution gathers the data from users on its own, nor the investigation of extracted OSN data is included in the category of market and public opinion research purposes, which are separately regulated in §28 (BDSG).

In contrast to work with personal attributes itself, we are merely interested in whether a user has published a single information or not because of the objective to point out attributes that could be used to link OSN profiles. It is not of interest for our investigation to analyze which concrete value of an attribute a user has published. Therefore, we do not need to permanently store a user's name, his published attributes, and his OSN profile itself. This concrete data is only needed during the automated analysis of profiles. Keeping it transiently in main memory is sufficient for our purpose. Such short-term processing of personal data separates the handling of personal information by OSNs, which has been authorized by the user, and the processing of merely statistical data that do not allow any conclusions regarding personal identities. Such shortest possible time gaps in the sense of separation of two domains of jurisdiction are referred to as a "logical second" [Win00]. The time span during which data is temporarily kept in main memory can only be defined as a logical second if it can be ensured that it is not possible to access the raw data of OSN profiles by a natural person. This applies to console outputs as well as to permanently stored data. As a consequence, working with a software which extracts only statistical data out of OSN profiles is compliant on the basis of German law only if it can be proven that nobody is able to recover personal attributes or profile information with a distinct link to a natural person or an OSN profile at runtime of the analysis software as well as afterwards. Hence, all processing steps in the analysis software have to be executed automatically.

4 Assessment of technical measures for privacy enhancement

As it is emphasized in Section 1, the publication of personal data bears a lot of risks such as the possibility for third-parties to implement software that crawls through OSNs and gathers information by emulation of users' behavior, respectively the technical behavior of a browser. Such software-based crawlers are often hard to distinguish from human users by OSNs. Before Section 5 presents the results of our statistical analysis on users' extractable data inside OSNs and the linkability of OSN profiles, we quickly discuss countermeasures of OSNs that are implemented to avoid automated extractions of personal data.

It is obvious, the better a crawler emulates the behavior of a human being and a browser, the more difficult it is to detect such intruding software. Hence, crawlers use previously registered accounts and determine parameters such as initialization vectors and valid session identifiers to login into OSNs. Additionally, to stay logged in crawlers have implemented a handling for cookies. The data will be extracted via parsing of HTML pages and identifying the information by means of tags or keywords inside the HTML code. Such keywords are commonly static and self-explanatory.

OSNs implemented several countermeasures to thwart and prevent crawling. If an alleged browser tries to get responses more and faster than ordinary users would try, CAPTCHAs are presented before the requested content is replied. CAPTCHAs are pictures with distorted letters that a user has to type in to be able to carry on surfing. CAPTCHAs are a common possibility to disturb crawling attacks, but an OSN has to weigh the safeness of its users and the potentially decreasing usability with respect to the amount of occurrences of CAPTCHAs. Into the bargain, today's CAPTCHA-challenges are breakable by software, too [ZYL⁺10]. Apart from that, a fix delay between sent HTTP-requests is often sufficient sophistication to prevent the occurrence of CAPTCHAs. Since only a valid e-mail address and some not further validated personal information is needed to create an OSN account, any time a profile is yet blocked by a CAPTCHA the software is able to register a new account and re-login into the OSN. This is also effective if OSNs are blocking further browsing for 24 hours after a fix amount of HTTP-gets within a fix amount of time. However, as long as these measures can be circumvented by re-login with another account, it is only a low extra effort for attackers to implement a bypass inside their crawlers.

Beyond these countermeasures, switches between standard websites (www...) and sites for mobile devices support crawlers to avoid Java Script (and its variations like AJAX) that are more difficult to emulate. Usually, the mobile pages that any common OSN provide are solely based on native HTML code and if the software is logged in into one of these alternatives the access to the other one is also open. It is obvious that such and other countermeasures are not sufficient to completely avoid automated data extraction.

5 Profile correlations based on friends lists

In this section we present a statistical evaluation on numerous profile pages of two of the most popular social networks in Germany and also worldwide. We refer to these OSNs

as OSN 1 and OSN 2 to not motivate attackers to implement crawlers in the future. We analyzed about 15,000 profiles of OSN 1 and about 50,000 profiles of OSN 2 to identify attributes that are suitable to be used to link OSN profiles of a single person from different OSNs (for instance by third-parties or to show a user how linkable his OSN profiles are). This analysis is done on the basis of the requirements stated in Section 3 and ensures, therefore, compliance to the German law as well as do not jeopardize users' privacy. Before this section presents the results of the study, the analysis software as well as its measures to ensure a lawful investigation are highlighted in the following.

The analysis software searches for a string (first name and last name) in different OSNs, gets lists of profiles, analyzes the publicly available information and compares each detected profile of one OSN with each detected one of the other OSNs. Afterwards, statistical data is extracted and stored permanently. Since the specific user name the software is searching for is not decisive for the results and in light of the presented restrictions concerning the fact that nobody should have access to a searched name, these input parameters have to be automatically generated. We implemented a fully automated generation of search parameters which in turn randomly picks and concatenates first names and last names out of large lists of popular names. A further characteristic of the software is that these input parameters as well as all extracted profile data are discarded subsequent to the analysis such that the analyzed profiles are unidentifiable.

To summarize, our analysis software autonomously generates input parameters for search queries, extracts publicly available information out of OSN profiles that are found with the searching term in different OSNs, calculates statistical data, tries to link profiles and discards all information that could identify the analyzed profiles. For sure, the requirement that the network will not be disturbed in terms of network load regarding numerous HTTP-requests by the developed software has additionally been taken into consideration.

In OSN 1 we figured out that 62.85% of profiles contain a publicly available birth date, whereas in OSN 2 only 1.24% of all analyzed profiles provide this information. In general, a users name and his birth date are applicable to identify a person out of a large circle of persons. Otherwise, it cannot be guaranteed that solely these two attributes appear uniquely in a group of people. Moreover, the probability that a combination of name and birth date is unique inside a large social network is obviously low and not sufficient to identify people. However, if it is an aim to link profiles of a specific person, in many cases these attributes provide enough information to establish such a link, depending on the amount of people who use the same name inside the networks.

In analogy to the birth date-attribute, some other attributes can be used to identify a user in an OSN. For instance, the name of the attended college or university is available in 53.88% of all analyzed profiles of OSN 1, in OSN 2 this information can be seen on 11.75% of profiles. A relationship status is provided by 25.56% (OSN 1), respectively 17.67% (OSN2) users. The hometown is available in 23.79% (OSN 1) or 13.36% (OSN 2) of the profiles.

Less suitable to link OSN profiles are attributes like the gender that is publicly available at 69.64% (OSN 1) or 71.31% (OSN 2) of the profiles. Since in the majority of cases the name implies the gender, this information has no additional value for the aim of user identification. However, 46.13% of profiles at OSN 1 provide a fully accessible friends

list and even 58.76% at OSN 2. In the following, we discuss the thesis that friends lists provides sufficient information to link OSN profiles.

This thesis is based on the assumption that a single combination of names that a friends list contain is to the greatest possible extent unique inside the network and, therefore, like a fingerprint of a single user. Furthermore, if such a single user is member of different OSNs, we assume that his friends lists overlap. Of course, if OSNs serve a particular purpose, i.e. for instance, business networking, the overlap of friends lists of OSNs that are dedicated for connecting to friends and family members are accordingly less. However, if this overlap is significantly higher than what we can expect for two OSN profiles with the same name but from different people, friends lists are a suitable type of information that can be easily (mis-)used to link profiles.

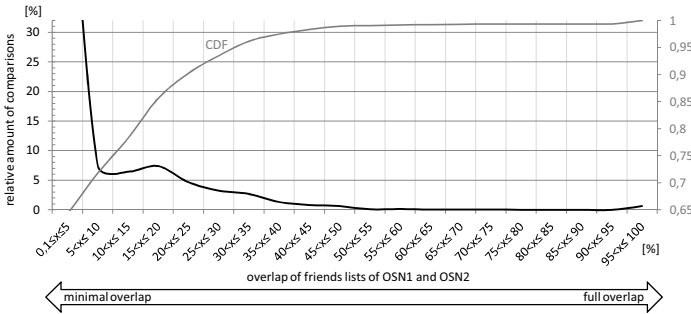


Figure 1: Histogram that shows the probability of friends lists overlap of two profiles from different OSNs with respect to the extent of the overlap and a corresponding CDF

Figure 1 shows the overlap of friends lists of profiles from two different OSNs. Therefore, we did more than 350,000 comparisons of friends lists, whereby in each comparison both analyzed profiles are registered under the same user name in the respective OSNs. The diagram shows the comparisons that result in overlaps between 0.1% and 100%. The x-axis shows the proportion of friends lists of OSN 1 that overlap with compared friends lists of OSN 2 and is divided into bins of 5%. On the y-axis the relative amount of comparisons within the different bins are drawn. The grey curve constitutes the corresponding cumulative distribution function with its y-axis on the right. It can be seen that the slope of the curve increases at about 7% overlap and has a local maximum at about 17% overlap.

The changed course at 7% overlap can be interpreted as a threshold which indicates that an identified overlap less than 7% could be a comparison of profiles of the same person or not. In general, the higher the overlap the higher the probability that the same person is owner of two compared profiles. Especially, this probability increases significantly if the comparison results in an overlap of more than 7%. With this interpretation, the local maximum constitutes an average overlap of friends lists of two profiles of a single person. However, the graph shows that a comparison of two friends lists of profiles of different OSNs results in many cases in no or just a minimal overlap. But, in some cases a higher overlap is obviously detectable and the results suggest that the overlap indicates a single person related to both compared profiles. Because of the requirements regarding compliance, it is not possible to validate such interpretation of the graphs. However, Table 1

Description	Example 1	Example 2	Example 3
Number of profiles in OSN 2 with identical name	111	193	97
Number of these profiles in OSN 2 with a friends list that overlap less than 2%	110	192	96
Number of these profiles in OSN 2 with a friends list that overlap more than 20%	1	1	1

Table 1: Examples of overlaps of single friends lists with friends lists from a different OSN

shows exemplarily the results of comparisons of three single profiles of OSN 1 with all profiles of OSN 2 wherein users registered with an identical name. It can be seen that in these examples exist a gap between overlaps less than 2% and overlaps higher than 20% which in turn underpins the interpretation of Figure 1.

6 Conclusions and future work

In OSNs users are faced with a thin line between publishing data to be part of the OSN as well as getting social recognitions such as it is shown in Section 1 and otherwise the potential of danger regarding privacy. In fact, third-parties are able to extract users' publicly available personal data without a permission of users and they are able to link information from different OSNs. In contrast to mandatory access control ensured by controlled and clear processes, e.g. as they are implemented in an identity management system, unintended data flows from OSNs to third-parties cannot be interrupted by users nor the OSN providers apart from adjusting privacy settings. The objective of the presented study is to demonstrate the linkability of OSN profiles if personal information is not hidden from strangers and therefore to show the possibilities of gathering data by third-parties.

In this paper, we firstly introduced a concept for statistical sampling on the basis of an analysis of German Data Protection Act that ensures compliance. Furthermore, we showed that the countermeasures OSNs have implemented to avoid unwanted data extraction by third-parties are not sufficient. Neither the integration of CAPTCHAs in page flows nor the blocking of further surfing after a fix amount of HTTP get-requests in a fix amount of time are effective to suppressing extractions entirely. In the following, we demonstrated that several attributes can be used to link information from different OSNs that belong to a single user. Moreover, we reinforced the thesis that friends lists, which are publicly accessible at about a half of all OSN profiles, can serve as fingerprints in OSNs and provide an opportunity to associate information. We substantiated the observation that an overlap of two compared friends lists of more than 7% is an indicator that both friends lists are part of profiles of a single natural person.

For future work, the presented study can be applied to constitute a basis for an application that shows users their concrete virtual appearance in a comprehensive manner. Particularly, the gained knowledge can be used to demonstrate a user the linkability of his OSN profiles and, thus, inadvertently exposed personal information can be unfolded. Such feedback

consequently results in an increased transparency that might motivate more users to adjust privacy settings in the future. An open issue to design such an application belongs to the question whether comparisons of a user's OSN profile with other profiles (matching the name of the profile holder) to determine overlaps are compliant to data protection acts because a processing of personal information of third OSN members cannot be avoided.

Acknowledgement

We like to thank Dr. iur. Oliver Raabe (KIT, ZAR) for his help in the analysis of German Data Protection Act and our students Irina Taranu, Florian Werling, and Michael Biebl for their contributions and help in the implementation of our analysis software.

References

- [GA05] R. Gross and A. Acquisti. Information revelation and privacy in online social networks. In *Proc. of the 2005 ACM workshop on Privacy in the electronic society*, WPES '05, pages 71–80, New York, NY, USA, 2005. ACM.
- [jim10] (german) JIM-STUDIE 2010 - Jugend, Information, (Multi-) Media. Technical report, Medienpaedagogischer Forschungsverbund Suedwest, Stuttgart, Germany, 2010.
- [KW08] B. Krishnamurthy and C. E. Wills. Characterizing privacy in online social networks. In *Proc. of the First Workshop on Online Social Networks*, WOSP '08, pages 37–42, New York, NY, USA, 2008. ACM.
- [KW10a] B. Krishnamurthy and C. E. Wills. On the leakage of personally identifiable information via online social networks. *SIGCOMM Comput. Commun. Rev.*, 2010.
- [KW10b] B. Krishnamurthy and C. E. Wills. Privacy leakage in mobile online social networks. In *Proc. of the 3rd conference on Online social networks*, WOSN'10, pages 4–4, Berkeley, CA, USA, 2010. USENIX Association.
- [LES07] C. A. Lampe, N. Ellison, and C. Steinfield. A familiar face(book): profile elements as signals in an online social network. In *Proc. of the SIGCHI conference on Human factors in computing systems*, CHI '07, pages 435–444, New York, USA, 2007. ACM.
- [Mis74] W. Mischel. Processes in Delay of Gratification. In *Advances in Experimental Social Psychology*, volume 7. Elsevier, 1974.
- [MKG⁺08] A. Mislove, H. S. Koppula, K. P. Gummadi, P. Druschel, and B. Bhattacharjee. Growth of the flickr social network. In *Proc. of the First Workshop on Online Social Networks*, WOSP '08, pages 25–30, New York, NY, USA, 2008. ACM.
- [MM83] H. N. Mischel and W. Mischel. The Development of Children's Knowledge of Self-Control Strategies. *Child Development*, 54(3):603–619, June 1983.
- [MV09] M. Motoyama and G. Varghese. I seek you: searching and matching individuals in social networks. In *Proceeding of the 11th International Workshop on Web Information and Data Management*, WIDM '09, pages 67–75, New York, USA, 2009. ACM.
- [TRW09] Mojtaba Torkjazi, Reza Rejaie, and Walter Willinger. Hot today, gone tomorrow: on the migration of MySpace users. In *Proc. of the 2nd ACM workshop on Online social networks*, WOSN '09, page 43–48, New York, NY, USA, 2009. ACM.
- [Win00] G. Winkler. (german) *Zeit und Recht. Kritische Anmerkungen zur Zeitgebundenheit des Rechts und des Rechtsdenkens*. Springer, 2000.
- [Zim04] G. Zimbardo. *Psychology and Life*. Pearson, 2004.
- [ZYL⁺10] B. B. Zhu, J. Yan, Q. Li, C. Yang, J. Liu, N. Xu, M. Yi, and K. Cai. Attacks and design of image recognition CAPTCHAs. In *Proc. of the 17th ACM Conference on Computer and Communications Security*, CCS '10, page 187–200, New York, USA, 2010. ACM.

IDMS und Datenschutz

Juristische Hürden und deren Überwindung

Mag. jur. Stephan Wagner, Dr. Marc Göcks
Projektleiter eCampus-IDMS
Multimedia Kontor Hamburg GmbH
Finkenau 31
D-22081 Hamburg
s.wagner@mmkh.de

Abstract: Bei der Durchführung von IT-Projekten stehen heutzutage oft weniger technische Probleme im Vordergrund. Vielmehr müssen organisatorische Rahmenbedingungen zur erfolgreichen Umsetzung geschaffen werden. Gerade im öffentlichen Bereich haben diese organisatorischen Bedingungen eine hohe juristische Komponente, wie z.B. Satzungen, Verordnungen und Gesetzestexte. Nachfolgend sollen die Komplexität und die spezifischen datenschutzrechtlichen Anforderungen für ein hochschulübergreifendes Identity Management am Beispiel des Hochschulstandortes Hamburg und seiner Spezifika im Rahmen des Hamburgischen Hochschul- und Datenschutzgesetz vorgestellt werden. Die Erfahrungen zeigen die Wichtigkeit der frühen Berücksichtigung juristischer Rahmenbedingungen und zeigen exemplarisch, wie diese geschaffen werden können.

1 Ausgangslage

Zeiten, an denen Studierende ihr Studium allein an ihrer Alma Mater absolvieren, gehören schon lange der Vergangenheit an. Praktika, Auslandssemester, praktische Jahre sind ebenso längst Bestandteil des Studierendenalltags geworden, wie die Literaturrecherche in den Bibliotheken benachbarter Hochschulen. Die Mobilität der Studierenden hat in den vergangenen Jahren und Jahrzehnten erheblich zugenommen. Entsprechendes gilt für die Mobilität der Wissenschaftlerinnen und Wissenschaftler an den Hochschulen. Diese Veränderungen führten an den Bibliotheken, Rechenzentren und vergleichbaren Hochschuleinrichtungen auch zu Änderungen im Nutzerverhalten, das wiederum zu geänderten Anforderungen an die Nutzerverwaltung dieser Einrichtungen und der von ihnen betriebenen IT-Verfahren geführt hat.

Studierende, wissenschaftliche Mitarbeiterinnen und Mitarbeiter, Professorinnen und Professoren und andere Angehörige des Personals der Hochschulen treten in verschiedenen Einrichtungen als Nutzer auf. So ist beispielsweise eine Studierende der Informatik an der Universität Hamburg nicht nur beim Studierendensekretariat der zentralen Universitätsverwaltung und im Campus-Management-System registriert, sondern muss sich auch für die Nutzung der Staats- und Universitätsbibliothek Hamburg (SUB), der Fachbereichsbibliothek und des Regionalen Rechenzentrums (RRZ) jeweils gesondert anmelden. Zieht sie um, so muss sie die Anschriftenänderung all diesen Einrichtungen jeweils gesondert mitteilen. Die Erfahrungen der Praxis zeigen, dass dies häufig nicht oder nur mangelhaft funktioniert. Die bei den verschiedenen Einrichtungen gespeicherten Datensätze weichen oftmals voneinander ab oder sind nicht auf dem aktuellen Stand.

An der (korrekten) Eintragung in solchen Verwaltungssystemen hängen oftmals wichtige Berechtigungen. Neben Berechtigungen zur Teilnahme an Lehrveranstaltungen oder Prüfungen sind hier insbesondere Berechtigungen zur Nutzung von IT-Systemen zu nennen, aber auch Nutzungs- oder Ausleihrechte für diverse Ressourcen wie Bibliotheken, Software, Mensen-Benutzung, etc.

Neben Sicherheitsüberlegungen, wie dem Schutz der persönlichen Daten und dem Schutz vor unberechtigter Nutzung von Ressourcen, kommen daher auch lizenzrechtliche Überlegungen ins Spiel.

Der Fall des Hochschulstandortes Hamburg erhält besondere Brisanz dadurch, dass hier hochschulübergreifende Kooperationen (die es seit jeher gibt) nun zeitgemäß durch vernünftige IT-Unterstützung widergespiegelt werden sollen.

2 Projektkontext und Rahmenbedingungen

Auf Basis dieser Ausgangslage und des u.a. damit verbundenen Modernisierungsdrucks sehen sich Hochschulinstitutionen und auch ganze Hochschulstandorte seit einigen Jahren der Herausforderung ausgesetzt, stärker die eigenen Prozesse und Services in der Lehre, der Forschung, dem Management und der Verwaltung zu analysieren und nach geeigneten Lösungen – vor allem auch IT-basierten Lösungen – zu suchen, welche zu einer Qualitäts- und auch Effizienzsteigerung beitragen können. Diesen sich stark verändernden Anforderungen hat sich der Hochschulstandort Hamburg schon vor einigen Jahren gestellt und ein gemeinsames, hochschulübergreifendes Projektprogramm „eCampus“ zur IT-basierten Modernisierung von Infrastrukturen sowie Verwaltungs- und Managementprozessen aufgesetzt. In Zusammenarbeit der sechs öffentlichen-staatlichen Hamburger Hochschulen (Universität Hamburg (UHH); Technische Universität Hamburg (TU); Hochschule für angewandte Wissenschaften (HAW); Hafen City Universität (HCU); Hochschule für Musik und Theater (HfMT); Hochschule für Bildende Künste (HfBK)) mit dem Multimedia Kontor Hamburg GmbH (MMKH) und der Behörde für Wissenschaft und Forschung werden gezielt IT-gestützte Modernisierungsprojekte an und für die Hochschulen initiiert.

In der erfolgreich durchgeführten ersten Projektphase von 2004-2006 lag der Fokus zunächst in der Identifikation von relevanten Anforderungs- und Problembereichen als auch im Aufbau von nachhaltigen Projektstrukturen. Diese umfassten sowohl die Bildung einer zentralen Lenkungsgruppe, als auch die Etablierung von themenspezifischen Arbeitsgruppen [HSV05]. Im Rahmen der Arbeitsgruppen sollten aktuelle Anforderungsbereiche für die Hamburger Hochschulen diskutiert, Erfahrungen ausgetauscht und erste exemplarische bzw. konzeptionelle Lösungsansätze erarbeitet werden. Auf der Grundlage der in den Arbeitsgruppen erzielten Ergebnisse und gewonnenen Erfahrungen sowie der aufgebauten Projektstrukturen wurde eCampus mit ausgewählten, priorisierten und zum Teil neu definierten Themenbereichen in weitere Projektierungsphasen überführt, welche aktuell eine Laufzeit bis Ende 2013 vorsehen. Innerhalb dieser Laufzeit sollen die einzelnen Aktivitäten in einen betriebsreifen Zustand überführt bzw. sollen die zum Teil bereits schon aufgebauten Produktiv- und Betriebsstrukturen weiter und nachhaltig etabliert werden.

Nachfolgend soll der Fokus vor allem auf das Teilprojekt für ein hochschulübergreifendes Identity Management System gelegt und dessen datenschutzrechtliche Besonderheiten herausgearbeitet werden.

3 eCampus-IDM – „Hochschulübergreifendes IDM“ am Standort Hamburg

Wie zuvor beschrieben, soll im Rahmen des eCampus Projektprogramms die Einführung eines hochschulübergreifenden Identity Management für die Hamburger Hochschulen realisiert werden. Die Besonderheit dieses Projektes liegt vor allem in seinem hochschulübergreifenden Ansatz und in der damit verbundenen großen Anzahl an beteiligten Partnerinstitutionen und zu konsolidierenden Identitäten. So studieren ca. 60.000 Studierende und arbeiten weit mehr als 15.000 Personen allein an den sechs öffentlich-staatlichen Hamburger Hochschulen [DG08], [Gö10].

3.1 Zielsetzung des Projektes

Ziel ist es, der selbstverständlichen Nutzung hochschulübergreifender Ressourcen, wie beispielsweise der Literaturrecherche in Bibliotheken benachbarter Hochschulen, für elektronische Services einen zeitgemäßen Rahmen zu geben. Dreh- und Angelpunkt bei Umsetzung dieser Anforderung ist ein intelligentes, hochschulübergreifendes Management personenbezogener Daten aller Hochschulmitglieder am Standort Hamburg.

Das eCampus IDMS soll alle Studierenden der verschiedenen Hochschulen derart konsolidieren, dass deren ursprüngliche Kennung über alle Hochschulen hinweg genutzt werden kann. Gleiches gilt letztendlich auch für Mitarbeiter und Lehrende.

Diese geänderten Anforderungen führen zu einem Bedarf nach einem einfachen, sicheren und komfortablen Zugriff auf hochschulübergreifende Daten. Dem wird die bisherige Situation der IT-Infrastruktur am Hochschulstandort Hamburg nicht gerecht. Dies soll sich mit der Einführung des eCampus-IDMS – ein intelligentes, hochschulübergreifendes Management personenbezogener Daten der Hochschulmitglieder am Standort Hamburg – nachhaltig verbessern [DG08]:

- Benötigte personenbezogene Daten der Hochschulmitglieder werden in ein zentrales System durch unterschiedliche Quellsysteme (Verwaltungssysteme der jeweiligen Hochschulen bzw. der SUB) eingespeist.
- Im zentralen System werden alle Daten automatisch konsolidiert und um überflüssige Informationen bereinigt.
- Nach Datenaufbereitung stellt das zentrale System aktuelle Mitgliedsdaten den dezentralen Systemen (IDMS der jeweiligen Hochschule bzw. der SUB) zur Verfügung.

Durch den oben beschriebenen Prozess der Daten-Konsolidierung wird erreicht:

- Doppelt eingespeiste Daten müssen nur einmal vorgehalten werden (Abstellung redundanter Datenspeicherung).
- Unterschiedliche Rollen einer Person (z.B. Bediensteter und Studierender in Personalunion) können der einen betroffenen Person zugeordnet werden (Schaffung einer elektronischen Identität, die grundlegend für die Umsetzung eines differenzierten Zugriffsberechtigungskonzepts ist).
- Nur die aktuellen und damit einzig gültigen personenbezogenen Daten werden gespeichert.
- Durch Vergabe einer hamburgweit eindeutigen Benutzer-Kennung wird erstmalig die institutionenübergreifende datenschutzrechtlich konforme Nutzung von elektronischen Services möglich (zentrale Freigabe und Sperrung von Zugriffsberechtigungen auf alle hochschulrelevanten Ressourcen und damit Vermeidung von „Datenleichen“).
- Für die Nutzung kostenpflichtiger Inhalte (z.B. Online-Magazine) kann sichergestellt werden, dass ehemalige Berechtigte nach ihrem Hochschulaustritt automatisch ihre Zugriffsberechtigung verlieren und somit nicht mehr unberechtigt Lizenzen „blockieren“.
- Es wird sichergestellt, dass Studierende unter einer einheitlichen und verifizierten elektronischen Identität die Dienste der Hamburger Hochschulen und der SUB nutzen (Vermeidung von Mehrfachkennungen und Mehrfachverwaltung).

Zusammenfassend ist das Ziel des hochschulübergreifenden Identity Management Systems, einen konsolidierten Bestand der Studierenden- und Mitarbeiter-Identitäten der Hamburger Hochschulen – ausschließlich zur Identifizierung und Autorisierung in IT-Anwendungsverfahren – bereitzustellen, grundlegende Verfahren für das Management dieser Identitäten zu etablieren und dieses System, soweit möglich, in Teilbereichen auch schon operativ verfügbar zu machen.

3.2 Architektur des eCampus-IDM

Die lokalen IDM-Systeme werden durch ein hochschulübergreifendes, zentrales Verzeichnis („Metadirectory eCampus-IDMS“) mit konsolidierten Identitäten versorgt. Diese Identitäten werden vorab von den lokalen Quell-Systemen der einzelnen Hochschulen an das eCampus-IDMS geliefert [GGW07]. Innerhalb des eCampus-IDMS ist dann eine Konsolidierung über alle Identitäten der beteiligten Hochschulen möglich, mithin die Etablierung von eindeutigen Personen-Identitäten. Diese konsolidierten Daten werden dann wieder an lokale IDM-Systeme (Ziel-Systeme) der Hochschulen provisioniert. Es wird damit sichergestellt, dass die lokalen IDM-Systeme mit dem gleichen konsolidierten Datenbestand arbeiten können und eine jeweils lokale Benutzerverwaltung über alle beteiligten Hochschul-Identitäten etabliert werden kann.

4 Problemstellung

Der oben skizzierte Ansatz eines hochschulübergreifenden IDM-Systems stellt unter anderem hohe Anforderungen an die zugrunde liegende Technik, sowie an ein gut strukturiertes Projektmanagement. Bei solchen Projekten wird oft der Fehler begangen, sie als reine IT-Projekte mit rein technischer Sichtweise umzusetzen. Die technische Umsetzung ist zwar durchaus eine sehr anspruchsvolle Aufgabe, kann aber bei gut durchgeführtem Projektmanagement in einem fest definierten Zeitfenster umgesetzt werden. Auch die Einbindung von den vielen verschiedenen Beteiligten aus den einzelnen Hochschulen stellt hohe Anforderungen an das eingesetzte Projektmanagement, jedoch ist dieser Part ebenfalls noch als überschaubar anzusehen.

Als ein viel größerer Problempunkt wurde die datenschutzrechtliche Freigabe eines hochschulübergreifenden IDM-Systems identifiziert. Sollen für ein hochschulübergreifendes Identitätsmanagement (IDMS) personenbezogene Daten aller Studierenden und aller Mitarbeiter von allen staatlichen Hochschulen Hamburgs zusammengeführt werden, so bedarf es hierfür einer rechtlichen Grundlage im Hamburgischen Hochschulgesetz (HmbHG). Die Daten sind auf das für die Aufgabenerfüllung erforderliche Maß zu beschränken. Trotz der geringen Sensibilität der einzelnen Daten sind angesichts der darüber erreichbaren Anwendungen Vorkehrungen für hohen bis sehr hohen Schutzbedarf zu treffen.

Personenbezogene Daten dürfen daher nur dann verarbeitet werden, wenn sie zur Aufgabenerfüllung erforderlich sind oder eine spezielle Vorschrift über den Datenschutz dies erlaubt. Die dritte Möglichkeit einer Einwilligungslösung scheidet von vornherein aus, da bei der unabdingbaren Möglichkeit eines jederzeitigen Widerrufs der Einwilligungserklärung das Verfahren unpraktikabel wäre und mithin eine Konsolidierung der Identitäten unnötig erschwert werden würde. Das HmbHG regelte ausdrücklich bisher lediglich die Datenverarbeitungsbefugnis der eigenen Hochschule gegenüber ihren Mitgliedern und würde somit nur die Einrichtung eines hochschuleigenen IDMS zulassen. Die allgemeinen Verarbeitungsvorschriften des HmbDSG setzen zudem voraus, dass die Erforderlichkeit zum Zeitpunkt der Verarbeitung besteht und kein milderes Mittel zur Verfügung steht.

Die Einrichtung eines hochschulübergreifenden IDM erforderte daher die Schaffung von Rechtsgrundlagen sowohl bezüglich des Datenaustausches als auch bezüglich seiner Automatisierung.

4.1 Fehlende Rechtsgrundlage

Mithin wurde ersichtlich, dass für die Etablierung eines hochschulübergreifenden IDM-Systems keine materiell-rechtliche Rechtsgrundlage vorlag. In der bisherigen Form des HmbHG konnte keine datenschutzrechtliche Freigabe konstruiert werden, da hier nur die einzelnen Hochschulen separat behandelt wurden und nicht in einem hochschulübergreifenden Kontext.

Im Folgenden wird der bisherige Wortlaut des § 111 HmbHG (Personenbezogene Daten) wiedergeben:

(1) Die Hochschulen dürfen von Studienbewerberinnen und Studienbewerbern, Studierenden, Prüfungskandidatinnen und Prüfungskandidaten, Absolventinnen und Absolventen sowie sonstigen Nutzerinnen und Nutzern von Hochschuleinrichtungen diejenigen personenbezogenen Daten erheben und verarbeiten, die für die Identifikation, die Zulassung, die Immatrikulation, die Rückmeldung, die Beurlaubung, die Teilnahme an Lehrveranstaltungen, die Prüfungen, die Nutzung von Hochschuleinrichtungen, die Hochschulplanung und die Kontaktpflege mit ehemaligen Hochschulmitgliedern erforderlich sind. Für Studierende kann zu diesem Zweck ein maschinenlesbarer Studierendenausweis eingeführt werden.

(2 - 3) ...

4.2 Darstellung der Problematik

Das bisherige Gesetz geht davon aus, dass der Einsatz von IDM-Systemen höchstens innerhalb einer Hochschule etabliert werden kann. Jedoch wäre mit diesem Ansatz nur die Basis für einen hochschulübergreifenden Ansatz geschaffen. Eine Konsolidierung aller Identitäten der Hamburger Hochschulen wäre nicht umzusetzen.

Auch das Umschwenken auf andere Ansätze erscheint vor dem Hintergrund der bisherigen Gesetzeslage schwierig, wenn diese die hochschulübergreifende Zusammenarbeit vernünftig widerspiegeln sollen und nicht wieder zu Insel-Lösungen ohne Datentransfer zwischen den einzelnen Hochschulen sowie der SUB führen soll. Ganz davon zu schweigen, dass ein solches Umschwenken zu einem späten Projektverlauf äußerst schwierig bis unmöglich ist.

4.3 Mögliche Folgen für das eCampus-IDMS

Aufgrund der vorliegenden Gesetzeslage wäre eine Umsetzung des Projektes nicht mehr zu gewährleisten, ein kompletter Stopp sämtlicher Aktivitäten wäre die Folge. Die bisher eingebrachten Ressourcen und Personalmittel wären unwiderruflich verloren.

5 Novellierung des HmbHG

Eine Novellierung des HmbHG war mithin die einzige Option, um das Projekt weiter führen zu können. In einem langen Prozess unter starker Einbindung des Hamburgischen Beauftragten für Datenschutz und Informationsfreiheit (HmbBfDI) [BfDI11] und des ansässigen Gesetzgebers konnte schließlich die folgende Neuformulierung des Hamburger Hochschulgesetzes erreicht werden:

Im Folgenden wird die neue Version des § 111 HmbHG (Personenbezogene Daten) wiedergeben:

(1 - 3) Gleichgeblieben

(4) Die Hochschulen und die Staats- und Universitätsbibliothek dürfen personenbezogene Daten ihrer Mitglieder, Angehörigen, Nutzerinnen und Nutzer im automatisierten Verfahren zusammenführen, soweit dies erforderlich ist zur Erstellung einer einheitlichen Benutzerkennung für von ihnen betriebene automatisierte Verfahren, Aktualisierung von Kommunikations- und Statusangaben sowie Einrichtung eines Kontos für elektronische Post.

Zu diesen Zwecken können sie eine gemeinsame automatisierte Datei einrichten. Eine Nutzung der Daten zu anderen Zwecken ist ausgeschlossen.

(5) Die Hochschulen regeln das Nähere durch Satzung, insbesondere welche Daten nach Absatz 1 erhoben und verarbeitet werden dürfen, die Aufbewahrungsfrist und das Verfahren bei der Ausübung des Auskunfts- und Einsichtsrechts, welche dieser Daten für die Zwecke der Hochschulstatistik verwendet und der dafür zuständigen Behörde übermittelt werden dürfen, die Daten und Funktionen eines maschinenlesbaren Studierendenausweises, die in diesem Zusammenhang nötigen Verfahrensregelungen sowie die Daten, die zur Erteilung des Ausweises erhoben und verarbeitet werden müssen, welche Daten nach Absatz 3 erhoben werden dürfen, die Verfahren der Erhebung dieser Daten sowie ihrer Verarbeitung und Auswertung, welche Daten nach Absatz 4 Satz 1 zusammengeführt werden dürfen und wie die gemeinsame Datei nach Absatz 4 Satz 2 auszugestalten ist. § 11a Absatz 1 Sätze 3 und 4 des Hamburgischen Datenschutzgesetzes gilt entsprechend. Betroffene können sich zur Wahrnehmung ihrer Rechte auf Auskunft, Berichtigung, Sperrung und Löschung an jede der beteiligten Stellen wenden.

(6) Der Senat wird ermächtigt, für den Bereich der Staats- und Universitätsbibliothek Regelungen nach Absatz 5 Nummer 5 durch Rechtsverordnung zu erlassen.

5.1 Darstellung der Novellierung

Im direkten Vergleich zwischen den beiden Versionen der Gesetze wird ersichtlich, dass nicht nur eine reine Anpassung des bisherigen Wortlauts vorgenommen wurde, sondern vielmehr ein komplett neuer Absatz eingearbeitet wurde.

5.2 Ablauf der Novellierung

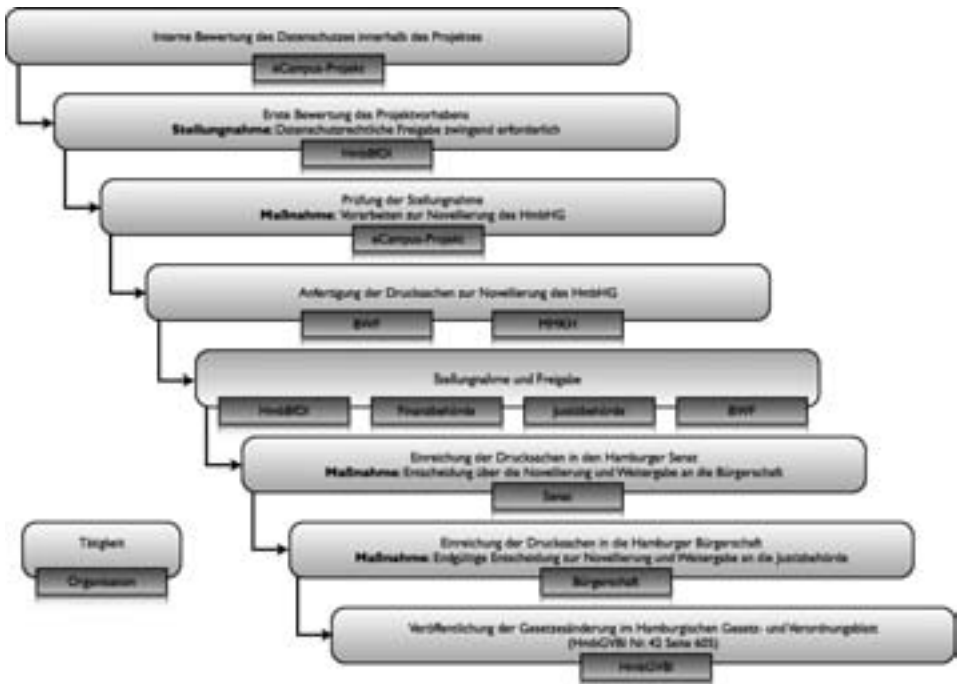


Abbildung 1: Ablaufdiagramm Novellierung

Innerhalb des Ablaufdiagramms wird schnell ersichtlich, dass an der Bearbeitung der Novellierung eine Vielzahl von unterschiedlichen Stellen, Behörden und Institutionen beteiligt waren. Wobei die federführende Koordination durch das MMKH übernommen wurde, welches beratend den verschiedenen Organisationseinheiten zur Verfügung stand.

5.3 Datenschutzrechtliche Anforderungen

Da die bisherigen Datenverarbeitungsregelungen entsprechend dem Erforderlichkeitsgrundsatz nur so weit gegriffen haben, wie die Datenverarbeitung zwischen der einzelnen Hochschule und ihren eigenen Mitgliedern geregelt ist (§ 111 HmbHG) und hochschulübergreifende Datenverarbeitung anhand der datenschutzrechtlichen Generalklauseln (§§ 12 ff HmbDSG) nur in den Einzelfällen erfolgen kann, in denen die Erforderlichkeit dazu zum Zeitpunkt der Verarbeitung gegeben ist, bedarf es einer gesonderten Befugnisnorm, wenn als Identitätsmanagementsystem ein hochschulübergreifendes Verfahren eingerichtet werden soll.

Darin ist das Metaverfahren mit dem kaskadierenden System rechtlich abzubilden. Wegen des mit der Verarbeitung verbundenen Eingriffs in das informationelle Selbstbestimmungsrecht der Betroffenen sind insbesondere die Verfahrensgrenzen und die Verantwortlichkeiten der einzelnen Beteiligten hinreichend bestimmt und normenklar abzubilden.

Bei dem Vorhaben, dass alle beteiligten Stellen die näher zu bestimmenden Daten elektronisch zusammenführen und in laufend konsolidierter Fassung aus dem Pool erhalten und für ihre Zwecke weiter verarbeiten, handelt es sich um eine gemeinsame Datei. Die Anforderungen des § 11 a HmbDSG sind daher entsprechend zu regeln. Einer zusätzlichen Ermächtigung für automatisierte Abrufe im Sinne des § 11 HmbDSG ist dafür nicht erforderlich.

Gleichzeitig ist sicherzustellen, dass ein für alle Beteiligten einheitliches und hinreichend sicheres Verfahren etabliert wird.

6 Folgen der Novellierung

Die Gesetzesänderungen führen nur dazu, dass der Austausch von Daten zwischen den Hochschulen untereinander und mit der SUB erleichtert wird. Dies erfolgt aber streng zweckgebunden: Der Datenaustausch darf nur dem Zweck dienen, eine einheitliche Nutzerverwaltung für die verschiedenen IT-Fachverfahren zu etablieren. Es ist nicht zulässig, Daten, die zur Nutzerverwaltung nicht erforderlich sind, auszutauschen. Insbesondere wird es mit diesem Gesetzentwurf nicht ermöglicht, fachliche Daten aus den IT-Fachverfahren ohne Bezug zum IDM auszutauschen. So wird es zwar ermöglicht, aus dem Datenbestand des Ausleihsystems der SUB die aktuelle Anschrift zu übermitteln; welche Bücher ausgeliehen wurden, darf aber nicht mitgeteilt werden. Daher droht nicht die Gefahr eines „gläsernen Wissenschaftlers“ oder Studierenden.

Über diese enge rechtliche Zweckbindung der Daten hinaus wird der Datenschutz auch durch organisatorisch-technische Maßnahmen geschützt, insbesondere durch die in § 8 HmbDSG benannten Maßnahmen wie z.B. die Zutrittskontrolle, die Zugangskontrolle, die Zugriffskontrolle, die Weitergabekontrolle, die Eingabekontrolle, die Verfügbarkeitskontrolle und das Gebot der Datentrennung.

7 Zusammenfassung

Die Erfahrungen im Hamburger Fall zeigen, dass eine rechtzeitige Berücksichtigung rechtlicher und organisatorischer Rahmenbedingungen nicht nur wichtig, sondern zwingend für eine erfolgreiche Umsetzung von IT-Projekten - insbesondere solchen, bei denen persönliche Daten verarbeitet werden - ist. Auch wenn dies innerhalb des Projekts durchaus bewusst war und Datenschützer sowie Gesetzgeber frühzeitig eingebunden wurden, hat es doch eine beträchtliche Zeit bis zur Umsetzung der Rechtsgrundlagen gebraucht und die Projektumsetzung verzögert. Hierbei gab es durchaus Momente, in denen eine erfolgreiche Umsetzung gar als komplett gefährdet angesehen wurde.

Trotz der Ähnlichkeit zu anderen Projekten, in denen organisationsübergreifende Verarbeitung von persönlichen Daten vorgesehen ist (wie die Shibboleth-Bestrebungen des DFN [DFN11], [Sh11] bzw. der Schweiz [Sw11] oder IntegraTUM [Inte06]) waren doch erhebliche Anstrengungen nötig, um eine Lösung für den Hochschulstandort Hamburg zu finden. Dies liegt unter anderem auch daran, dass die gesetzlichen Bedingungen in den verschiedenen Bundesländern sehr unterschiedlich sein können, und daran, dass die konkrete Umsetzung des Hamburger Projekts anders als in anderen Projekten aussieht.

Mit der Gesetzesänderung ist ein erheblicher Meilenstein bei der Umsetzung von eCampus-IDMS erreicht. Auch wenn durchaus weitere organisatorische Bedingungen es nicht gänzlich erlauben, sich nun komplett auf eine “nur” technische Umsetzung zu stürzen, ist damit eine wesentliche Hürde zur erfolgreichen Umsetzung genommen.

Literaturverzeichnis

- [DFN11] Homepage der DFN-AAI; <https://www.aai.dfn.de/>
- [DG08] Dietz, G.; Göcks, M.: Serviceorientierung durch hochschulübergreifendes Identitätsmanagement. In (Hegering, H-G.; Lehmann, A.; Ohlbach, H. J.; Scheideler, C., Hrsg.): INFORMATIK 2008 Beherrschbare Systeme – dank Informatik Band 2. GI Jahrestagung 2008, Bonn, 2008, S. 577-582.
- [GGW07] Gennis, M.; Gradmann, S.; Winklmeier, S.: IDM@eCampus.HH – Identitätsmanagement an Hamburger Hochschulen. In (Gaedke, M.; Borgeest, R., Hrsg.): Integriertes Informationsmanagement an Hochschulen – Quo Vadis Universität 2.0?, Karlsruhe, Universitätsverlag Karlsruhe, 2007, S. 91-107.
- [Gö10] Göcks, M.: Hochschulübergreifende Service- und Beratungsstrukturen des Wissenschaftsstandortes Hamburg am Beispiel des MMKH. In (Bremer, C.; Göcks, M.; Rühl, P.; Stratmann, J., Hrsg.): Landesinitiativen für E-Learning an deutschen Hochschulen, Waxmann Verlag, 2010,
- [HSV05] Haussner, S.; Schmid, U.; Vogel, M.: Vom e-Learning zum eCampus. Hamburgs Hochschulen auf dem Weg zu einer integrierten e-Learning- und IT-Dienste-Infrastruktur. In: Zeitschrift für Hochschuldidaktik (ZFHD), April 2005, S. 33-46; http://www.zfhd.at/resources/downloads/→ZFHD_03_03_Haussner_eCampus_HH_1000343.pdf
- [Inte06] IntegraTUM. Bericht 2004-2006, Antrag 2006-2009. online unter http://portal.mytum.de/iuk/integratum/dokumente/index_html/vortrag_DFG.pdf
- [Sh11] Homepage von Shibboleth; <http://shibboleth.internet2.edu/>
- [Sw11] Homepage von SWITCH AAI; <http://www.switch.ch/aai/>
- [BfDI11] Homepage des HmbBfDI; <http://www.datenschutz-hamburg.de/>

Kurzbeschreibung der Poster

EVALSO,
an enabling communication infrastructure
for
Astronomy

R.Lemke^{a)}, G.Filippi^{b)}, R.Chini^{a)}

a) Ruhr-Universität Bochum
Astronomisches Institut
Universitätsstraße 150
44801 Bochum
rlemke@astro.ruhr-uni-bochum.de
rchini@astro.ruhr-uni-bochum.de

b) European Southern Observatory (ESO)
Karl-Schwarzschild-Strasse 2
85748 Garching bei München
gfilippi@eso.org

Abstract: EVALSO (Enabling Virtual Access to Latin-American Southern Observatories) is an international consortium of nine astronomical organisations, and research network operators, part-funded under the European Commission FP7, to create and exploit high-speed bandwidth connections to South American observatories. The paper, after describing the genesis of the project and the selection and construction of the communication infrastructure, analyses the application aspects and the possibilities that the new infrastructure can open within the astronomical and education community, special attention is given to the experience done at RUB, one of the project members.

Dedizierte Wellenlängenverbindungen als reguläres Dienstangebot der europäischen Forschungsnetze

Andreas Hanemann^{*}, Wolfgang Fritz[#], Mark Yampolskiy[#]

^{*} DFN-Verein
Alexanderplatz 1, 10178 Berlin
hanemann@dfn.de

[#] Leibniz-Rechenzentrum München
Boltzmannstr. 1, 85748 Garching bei München
{wolfgang.fritz,mark.yampolskiy}@lrz.de

In den vergangenen Jahren wurden für einige Projekte Verbindungen mit dedizierten Wellenlängen über verschiedene Netze hinweg geschaltet, wobei das private optische Netz für die Large Hadron Collider-Experimente am CERN das prominenteste Beispiel darstellt. Diese Realisierungen beruhen jeweils auf speziellen Vereinbarungen. Um die Bereitstellung solcher Wellenlängenverbindungen in Zukunft zu vereinfachen, wird im Rahmen des GÉANT-Projektes eine allgemeine Dienstdefinition vorgenommen. Diese soll nach einer Prototypphase im Frühjahr/Sommer 2011 für zukünftige Verbindungen dieser Art verwendet werden.

Bei der Dienstdefinition spielen zwei Dokumente eine zentrale Rolle. Die sog. Service Level Specification (SLS) stellt die Schnittstelle zwischen Einrichtungen, die den Dienst nutzen wollen, und den GÉANT-Partnern dar, während die Zusammenarbeit zwischen den GÉANT-Partnern in einem Operational Level Agreement (OLA) geregelt wird. Im SLS werden die allgemeinen Eigenschaften des Dienstes beschrieben, beispielsweise welche Verfügbarkeit angestrebt wird und wie Nutzer Zugriff auf aktuelle Messdaten bekommen können. Im OLA wird spezifiziert, wie die Eigenschaften des Dienstes erreicht werden können. Hierbei ist insbesondere die Definition von Rollen und Prozessen wichtig sowie der Einsatz von Software Tools.

Zwei Software Tools sind für den Dienst besonders zu erwähnen. Das am Leibniz-Rechenzentrum entwickelte Monitoring-System *E2Emon* wird seit Ende 2006 für die kontinuierliche Überwachung von Verbindungen auf den ISO/OSI-Schichten 1 und 2 verwendet. Dazu werden relevante Informationen aus den Managementsystemen der beteiligten Forschungsnetze korreliert und ausgewertet. Der Informationsaustausch und die Koordination bei der Planung von diesen Verbindungen soll in naher Zukunft vom Tool *I-SHARe* (Information Sharing across Heterogeneous Administrative Regions) unterstützt werden. Für die Entwicklung des Tools wurden insbesondere vom italienischen Forschungsnetz GARR und vom DFN die Prozessabläufe bei der Zusammenarbeit der GÉANT-Partner analysiert, modelliert und schließlich in dem Tool abgebildet.

HVMcast: Entwicklung und Evaluierung einer Architektur zur universellen Gruppenkommunikation im Internet *

Sebastian Meiling¹, Dominik Charousset¹, Thomas C. Schmidt¹, Matthias Wählisch^{2,3}
{sebastian.meiling, dominik.charousset}@haw-hamburg.de; {t.schmidt, waehlich}@ieee.org

¹HAW Hamburg, Dept. Informatik, Berliner Tor 7, 20099 Hamburg

²FU Berlin, Institut für Informatik, Takustr. 9, 14195 Berlin

³link-lab, Hönow Str. 35, 10318 Berlin

Gruppenkommunikation ist die Basis vieler Internetanwendungen wie IPTV und Online Multiplayer Spiele. Trotz existierender Multicast-Verfahren für eine effiziente Datenverteilung, setzen diese Anwendungen oft auf proprietäre Technologien, weil bisher eine technologieübergreifende, transparente Multicastschnittstelle sowie ein allgemeiner Multicastdienst im Internet fehlen. In diesem Beitrag präsentieren wir die Implementierung und Evaluierung der HVMcast-Architektur [MCSW10] zur Bereitstellung eines universellen Multicast. HVMcast beruht auf einem abstrakten Namensschema nebst Anwendungsschnittstelle [WSV11], einer systemzentrierte Middleware für Endsysteme sowie Gateways zur Verknüpfung unterschiedlicher Technologien und Multicast-Domänen.

Das Grundkonzept von HVMcast ist ein *Late-Binding* von Gruppennamen zu technologie-spezifischen Gruppenadressen (Loc-ID Split) über eine allgemeine Multicast-API. Anwendungen können somit unabhängig von den zur Laufzeit verfügbaren Technologien entwickelt und kompiliert werden. Das Architekturkonzept wurde in Form einer Middleware-Komponente (Abb. 1) und einer API-Bibliothek für Gruppenanwendungen implementiert. Multicast-Technologien sind dabei als Module gekapselt und können entsprechend der lokalen Systemumgebung dynamisch geladen werden. In der Evaluierung wurde die Performance der Middleware im Vergleich zum Linux IP-Stack in Abhängigkeit zur Paketgröße analysiert. Gemessen wurden Paketdurchsatz und -verlust sowie die CPU-Auslastung beim Sender und Empfänger. Die Ergebnisse zeigen, dass die Performance der Middleware an den Paketdurchsatz (#/s) gekoppelt ist und für Paketgrößen ab 500 B eine 1 Gbit/s Netzwerkverbindung auslasten kann.

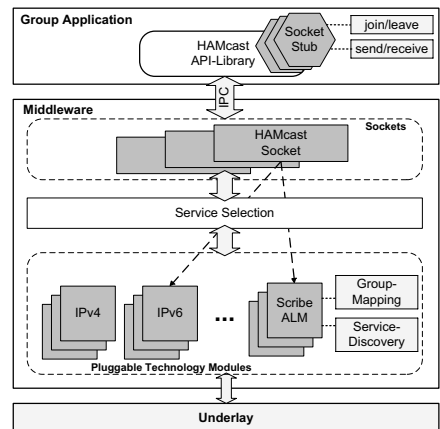


Abbildung 1: Aufbau der HVMcast-Middleware und API-Bibliothek.

Literatur

- [MCSW10] S. Meiling, D. Charousset, T. C. Schmidt und M. Wählisch. System-assisted Service Evolution for a Future Internet – The HAMcast Approach to Pervasive Multicast. In *Proc. of IEEE GLOBECOM'10, Workshop MCS*, IEEE Press, Dec. 2010.
- [WSV11] M. Waehlich, T. Schmidt und S. Venaas. A Common API for Transparent Hybrid Multicast. Internet-Draft – work in progress 01, IETF, March 2011.

*Diese Arbeit wird vom Bundesministerium für Bildung und Forschung im Rahmen des Projekts HVMcast in der G-Lab Initiative gefördert, s. <http://hamcast.realmv6.org>.

Drahtlose Sensornetze in der Lehre – Sicherheitsbetrachtung von ZigBee-Netzen

Björn Stelte
Institut für Technische Informatik
Universität der Bundeswehr München
bjoern.stelte@unibw.de

Gerade in den letzten Jahren kommen mehr und mehr Anwendungen auf Basis von drahtlosen Sensornetze auf dem Markt. So gibt es in den USA die Bestrebungen drahtlose Sensornetze für die Energieüberwachung (smart monitoring) einzusetzen, aber auch in der Hausautomatisierung, dem Gesundheitswesen, u.a. sind die Netze bestehend aus vielen kleinen Sensorknoten im Gespräch. Der de-facto Standard bzgl. der Kommunikation dieser Sensornetze ist derzeit ZigBee. Dieses Protokoll wird in nahezu allen gängigen Sensornetzen verwendet. Jedoch sind bereits seit einiger Zeit Schwächen dieses Protokolls bekannt. Wie Joshua Wright mit seinem KillerBee Framework gezeigt hat, ist ein praktischer Angriff auf eine ZigBee Infrastruktur mit wenigen Hilfsmitteln kostengünstig und schnell realisierbar. Dieser Beitrag gibt einen Überblick über das ZigBee Protokoll und einige seiner Schwachstellen, sowie über Werkzeuge, um die Problematik im Unterricht darzustellen.

In der Lehre ist es häufig erforderlich zum einen aktuelle und spannende Themen aufzugreifen aber zum anderen auch diese den Studenten geeignet zu vermitteln. Eine Demonstration von Kommunikationsabläufen sicherheitskritischer Datenkommunikation und insbesondere die Erklärung der Funktionsweise von Angriffen auf selbige anhand einer Simulation erhöht den Lerneffekt. Die in diesem Beitrag aufgezeigte Sicherheits-Problematik des ZigBee Protokolls ist ein aktuelles Thema. Die vorgeschlagenen Simulationen können bspw. im Unterricht oder bei Beteiligung der Studenten in einem Praktikum verwendet werden, um die Aspekte der IT-Sicherheit dem Studenten nahezubringen. Wir haben die Erfahrung gesammelt, dass der Drei-Satz Erkennen, Analysieren, Absichern den besten Lerneffekt erzielt. So sind die drei vorgestellten Simulationen in genau dieser Reihung zu sehen. Die erste Simulation schult das Erkennen der Problematik, die zweite Simulation vermittelt eine Analyse des Angriffs (hier KillerBee) und letztendlich wird in der dritten Simulation die Absicherung des Netzes anhand eines Beispiels erlernt. Weitere Angriffe auf ZigBee Netze sind möglich, so kann bei Interesse die Simulationsaufgabe nach belieben noch erweitert oder ausgebaut werden.

Der vollständige Beitrag kann unter der Adresse http://inf3-www.informatik.unibw-muenchen.de/ZigBee_sec_lehre.pdf abgerufen werden.

GI-Edition Lecture Notes in Informatics

- P-1 Gregor Engels, Andreas Oberweis, Albert Zündorf (Hrsg.): Modellierung 2001.
- P-2 Mikhail Godlevsky, Heinrich C. Mayr (Hrsg.): Information Systems Technology and its Applications, ISTA'2001.
- P-3 Ana M. Moreno, Reind P. van de Riet (Hrsg.): Applications of Natural Language to Information Systems, NLDB'2001.
- P-4 H. Wörn, J. Mühling, C. Vahl, H.-P. Meinzer (Hrsg.): Rechner- und sensorgestützte Chirurgie; Workshop des SFB 414.
- P-5 Andy Schürr (Hg.): OMER – Object-Oriented Modeling of Embedded Real-Time Systems.
- P-6 Hans-Jürgen Appelrath, Rolf Beyer, Uwe Marquardt, Heinrich C. Mayr, Claudia Steinberger (Hrsg.): Unternehmen Hochschule, UH'2001.
- P-7 Andy Evans, Robert France, Ana Moreira, Bernhard Rumpe (Hrsg.): Practical UML-Based Rigorous Development Methods – Countering or Integrating the extremists, pUML'2001.
- P-8 Reinhard Keil-Slawik, Johannes Magenheimer (Hrsg.): Informatikunterricht und Medienbildung, INFOS'2001.
- P-9 Jan von Knop, Wilhelm Haverkamp (Hrsg.): Innovative Anwendungen in Kommunikationsnetzen, 15. DFN Arbeitstagung.
- P-10 Mirjam Minor, Steffen Staab (Hrsg.): 1st German Workshop on Experience Management: Sharing Experiences about the Sharing Experience.
- P-11 Michael Weber, Frank Kargl (Hrsg.): Mobile Ad-Hoc Netzwerke, WMAN 2002.
- P-12 Martin Glinz, Günther Müller-Luschnat (Hrsg.): Modellierung 2002.
- P-13 Jan von Knop, Peter Schirmbacher and Viljan Mahni_ (Hrsg.): The Changing Universities – The Role of Technology.
- P-14 Robert Tolksdorf, Rainer Eckstein (Hrsg.): XML-Technologien für das Semantic Web – XSW 2002.
- P-15 Hans-Bernd Bludau, Andreas Koop (Hrsg.): Mobile Computing in Medicine.
- P-16 J. Felix Hampe, Gerhard Schwabe (Hrsg.): Mobile and Collaborative Business 2002.
- P-17 Jan von Knop, Wilhelm Haverkamp (Hrsg.): Zukunft der Netze –Die Verletzbarkeit meistern, 16. DFN Arbeitstagung.
- P-18 Elmar J. Sinz, Markus Plaha (Hrsg.): Modellierung betrieblicher Informationssysteme – MobIS 2002.
- P-19 Sigrid Schubert, Bernd Reusch, Norbert Jesse (Hrsg.): Informatik bewegt – Informatik 2002 – 32. Jahrestagung der Gesellschaft für Informatik e.V. (GI) 30.Sept.-3.Okt. 2002 in Dortmund.
- P-20 Sigrid Schubert, Bernd Reusch, Norbert Jesse (Hrsg.): Informatik bewegt – Informatik 2002 – 32. Jahrestagung der Gesellschaft für Informatik e.V. (GI) 30.Sept.-3.Okt. 2002 in Dortmund (Ergänzungsband).
- P-21 Jörg Desel, Mathias Weske (Hrsg.): Promise 2002: Prozessorientierte Methoden und Werkzeuge für die Entwicklung von Informationssystemen.
- P-22 Sigrid Schubert, Johannes Magenheimer, Peter Hubwieser, Torsten Brinda (Hrsg.): Forschungsbeiträge zur “Didaktik der Informatik” – Theorie, Praxis, Evaluation.
- P-23 Thorsten Spitta, Jens Borchers, Harry M. Sneed (Hrsg.): Software Management 2002 – Fortschritt durch Beständigkeit
- P-24 Rainer Eckstein, Robert Tolksdorf (Hrsg.): XMIDX 2003 – XML-Technologien für Middleware – Middleware für XML-Anwendungen
- P-25 Key Pousttchi, Klaus Turowski (Hrsg.): Mobile Commerce – Anwendungen und Perspektiven – 3. Workshop Mobile Commerce, Universität Augsburg, 04.02.2003
- P-26 Gerhard Weikum, Harald Schöning, Erhard Rahm (Hrsg.): BTW 2003: Datenbanksysteme für Business, Technologie und Web
- P-27 Michael Kroll, Hans-Gerd Lipinski, Kay Melzer (Hrsg.): Mobiles Computing in der Medizin
- P-28 Ulrich Reimer, Andreas Abecker, Steffen Staab, Gerd Stumme (Hrsg.): WM 2003: Professionelles Wissensmanagement – Erfahrungen und Visionen
- P-29 Antje Düsterhöft, Bernhard Thalheim (Eds.): NLDB'2003: Natural Language Processing and Information Systems
- P-30 Mikhail Godlevsky, Stephen Liddle, Heinrich C. Mayr (Eds.): Information Systems Technology and its Applications
- P-31 Arslan Brömme, Christoph Busch (Eds.): BIOSIG 2003: Biometrics and Electronic Signatures

- P-32 Peter Hubwieser (Hrsg.): Informatische Fachkonzepte im Unterricht – INFOS 2003
- P-33 Andreas Geyer-Schulz, Alfred Taudes (Hrsg.): Informationswirtschaft: Ein Sektor mit Zukunft
- P-34 Klaus Dittrich, Wolfgang König, Andreas Oberweis, Kai Rannenber, Wolfgang Wahlster (Hrsg.): Informatik 2003 – Innovative Informatikanwendungen (Band 1)
- P-35 Klaus Dittrich, Wolfgang König, Andreas Oberweis, Kai Rannenber, Wolfgang Wahlster (Hrsg.): Informatik 2003 – Innovative Informatikanwendungen (Band 2)
- P-36 Rüdiger Grimm, Hubert B. Keller, Kai Rannenber (Hrsg.): Informatik 2003 – Mit Sicherheit Informatik
- P-37 Arndt Bode, Jörg Desel, Sabine Rathmayer, Martin Wessner (Hrsg.): DeLFI 2003: e-Learning Fachtagung Informatik
- P-38 E.J. Sinz, M. Plaha, P. Neckel (Hrsg.): Modellierung betrieblicher Informationssysteme – MobIS 2003
- P-39 Jens Nedon, Sandra Frings, Oliver Göbel (Hrsg.): IT-Incident Management & IT-Forensics – IMF 2003
- P-40 Michael Rebstock (Hrsg.): Modellierung betrieblicher Informationssysteme – MobIS 2004
- P-41 Uwe Brinkschulte, Jürgen Becker, Dietmar Fey, Karl-Erwin Großpietsch, Christian Hochberger, Erik Maehle, Thomas Runkler (Edts.): ARCS 2004 – Organic and Pervasive Computing
- P-42 Key Pousttchi, Klaus Turowski (Hrsg.): Mobile Economy – Transaktionen und Prozesse, Anwendungen und Dienste
- P-43 Birgitta König-Ries, Michael Klein, Philipp Obreiter (Hrsg.): Persistence, Scalability, Transactions – Database Mechanisms for Mobile Applications
- P-44 Jan von Knop, Wilhelm Haverkamp, Eike Jessen (Hrsg.): Security, E-Learning, E-Services
- P-45 Bernhard Rumpe, Wolfgang Hesse (Hrsg.): Modellierung 2004
- P-46 Ulrich Flegel, Michael Meier (Hrsg.): Detection of Intrusions of Malware & Vulnerability Assessment
- P-47 Alexander Prosser, Robert Krimmer (Hrsg.): Electronic Voting in Europe – Technology, Law, Politics and Society
- P-48 Anatoly Doroshenko, Terry Halpin, Stephen W. Liddle, Heinrich C. Mayr (Hrsg.): Information Systems Technology and its Applications
- P-49 G. Schiefer, P. Wagner, M. Morgenstern, U. Rickert (Hrsg.): Integration und Datensicherheit – Anforderungen, Konflikte und Perspektiven
- P-50 Peter Dadam, Manfred Reichert (Hrsg.): INFORMATIK 2004 – Informatik verbindet (Band 1) Beiträge der 34. Jahrestagung der Gesellschaft für Informatik e.V. (GI), 20.-24. September 2004 in Ulm
- P-51 Peter Dadam, Manfred Reichert (Hrsg.): INFORMATIK 2004 – Informatik verbindet (Band 2) Beiträge der 34. Jahrestagung der Gesellschaft für Informatik e.V. (GI), 20.-24. September 2004 in Ulm
- P-52 Gregor Engels, Silke Seehusen (Hrsg.): DELFI 2004 – Tagungsband der 2. e-Learning Fachtagung Informatik
- P-53 Robert Giegerich, Jens Stoye (Hrsg.): German Conference on Bioinformatics – GCB 2004
- P-54 Jens Borchers, Ralf Kneuper (Hrsg.): Softwaremanagement 2004 – Outsourcing und Integration
- P-55 Jan von Knop, Wilhelm Haverkamp, Eike Jessen (Hrsg.): E-Science und Grid Ad-hoc-Netze Medienintegration
- P-56 Fernand Feltz, Andreas Oberweis, Benoit Otjacques (Hrsg.): EMISA 2004 – Informationssysteme im E-Business und E-Government
- P-57 Klaus Turowski (Hrsg.): Architekturen, Komponenten, Anwendungen
- P-58 Sami Beydeda, Volker Gruhn, Johannes Mayer, Ralf Reussner, Franz Schweiggert (Hrsg.): Testing of Component-Based Systems and Software Quality
- P-59 J. Felix Hampe, Franz Lehner, Key Pousttchi, Kai Rannenber, Klaus Turowski (Hrsg.): Mobile Business – Processes, Platforms, Payments
- P-60 Steffen Friedrich (Hrsg.): Unterrichtskonzepte für informatische Bildung
- P-61 Paul Müller, Reinhard Gotzhein, Jens B. Schmitt (Hrsg.): Kommunikation in verteilten Systemen
- P-62 Federrath, Hannes (Hrsg.): „Sicherheit 2005“ – Sicherheit – Schutz und Zuverlässigkeit
- P-63 Roland Kaschek, Heinrich C. Mayr, Stephen Liddle (Hrsg.): Information Systems – Technology and its Applications

- P-64 Peter Liggesmeyer, Klaus Pohl, Michael Goedicke (Hrsg.): Software Engineering 2005
- P-65 Gottfried Vossen, Frank Leymann, Peter Lockemann, Wolffried Stucky (Hrsg.): Datenbanksysteme in Business, Technologie und Web
- P-66 Jörg M. Haake, Ulrike Lucke, Djamshid Tavangarian (Hrsg.): DeLFI 2005: 3. deutsche e-Learning Fachtagung Informatik
- P-67 Armin B. Cremers, Rainer Manthey, Peter Martini, Volker Steinhage (Hrsg.): INFORMATIK 2005 – Informatik LIVE (Band 1)
- P-68 Armin B. Cremers, Rainer Manthey, Peter Martini, Volker Steinhage (Hrsg.): INFORMATIK 2005 – Informatik LIVE (Band 2)
- P-69 Robert Hirschfeld, Ryszard Kowalczyk, Andreas Polze, Matthias Weske (Hrsg.): NODe 2005, GSEM 2005
- P-70 Klaus Turowski, Johannes-Maria Zaha (Hrsg.): Component-oriented Enterprise Application (COAE 2005)
- P-71 Andrew Torda, Stefan Kurz, Matthias Rarey (Hrsg.): German Conference on Bioinformatics 2005
- P-72 Klaus P. Jantke, Klaus-Peter Fährnrich, Wolfgang S. Wittig (Hrsg.): Marktplatz Internet: Von e-Learning bis e-Payment
- P-73 Jan von Knop, Wilhelm Haverkamp, Eike Jessen (Hrsg.): "Heute schon das Morgen sehen"
- P-74 Christopher Wolf, Stefan Lucks, Po-Wah Yau (Hrsg.): WEWoRC 2005 – Western European Workshop on Research in Cryptology
- P-75 Jörg Desel, Ulrich Frank (Hrsg.): Enterprise Modelling and Information Systems Architecture
- P-76 Thomas Kirste, Birgitta König-Riess, Key Pousttchi, Klaus Turowski (Hrsg.): Mobile Informationssysteme – Potentiale, Hindernisse, Einsatz
- P-77 Jana Dittmann (Hrsg.): SICHERHEIT 2006
- P-78 K.-O. Wenkel, P. Wagner, M. Morgens-tern, K. Luzi, P. Eisermann (Hrsg.): Land- und Ernährungswirtschaft im Wandel
- P-79 Bettina Biel, Matthias Book, Volker Gruhn (Hrsg.): Softwareengineering 2006
- P-80 Mareike Schoop, Christian Huemer, Michael Rebstock, Martin Bichler (Hrsg.): Service-Oriented Electronic Commerce
- P-81 Wolfgang Karl, Jürgen Becker, Karl-Erwin Großpietsch, Christian Hochberger, Erik Maehle (Hrsg.): ARCS '06
- P-82 Heinrich C. Mayr, Ruth Breu (Hrsg.): Modellierung 2006
- P-83 Daniel Huson, Oliver Kohlbacher, Andrei Lupas, Kay Nieselt and Andreas Zell (eds.): German Conference on Bioinformatics
- P-84 Dimitris Karagiannis, Heinrich C. Mayr, (Hrsg.): Information Systems Technology and its Applications
- P-85 Witold Abramowicz, Heinrich C. Mayr, (Hrsg.): Business Information Systems
- P-86 Robert Krimmer (Ed.): Electronic Voting 2006
- P-87 Max Mühlhäuser, Guido Rößling, Ralf Steinmetz (Hrsg.): DELFI 2006: 4. e-Learning Fachtagung Informatik
- P-88 Robert Hirschfeld, Andreas Polze, Ryszard Kowalczyk (Hrsg.): NODe 2006, GSEM 2006
- P-90 Joachim Schelp, Robert Winter, Ulrich Frank, Bodo Rieger, Klaus Turowski (Hrsg.): Integration, Informationslogistik und Architektur
- P-91 Henrik Stormer, Andreas Meier, Michael Schumacher (Eds.): European Conference on eHealth 2006
- P-92 Fernand Feltz, Benoît Otjacques, Andreas Oberweis, Nicolas Poussing (Eds.): AIM 2006
- P-93 Christian Hochberger, Rüdiger Liskowsky (Eds.): INFORMATIK 2006 – Informatik für Menschen, Band 1
- P-94 Christian Hochberger, Rüdiger Liskowsky (Eds.): INFORMATIK 2006 – Informatik für Menschen, Band 2
- P-95 Matthias Weske, Markus Nüttgens (Eds.): EMISA 2005: Methoden, Konzepte und Technologien für die Entwicklung von dienstbasierten Informationssystemen
- P-96 Saartje Brockmans, Jürgen Jung, York Sure (Eds.): Meta-Modelling and Ontologies
- P-97 Oliver Göbel, Dirk Schadt, Sandra Frings, Hardo Hase, Detlef Günther, Jens Nedon (Eds.): IT-Incident Mangament & IT-Forensics – IMF 2006

- P-98 Hans Brandt-Pook, Werner Simonsmeier und Thorsten Spitta (Hrsg.): Beratung in der Softwareentwicklung – Modelle, Methoden, Best Practices
- P-99 Andreas Schwill, Carsten Schulte, Marco Thomas (Hrsg.): Didaktik der Informatik
- P-100 Peter Forbrig, Günter Siegel, Markus Schneider (Hrsg.): HDI 2006: Hochschuldidaktik der Informatik
- P-101 Stefan Böttinger, Ludwig Theuvsen, Susanne Rank, Marlies Morgenstern (Hrsg.): Agrarinformatik im Spannungsfeld zwischen Regionalisierung und globalen Wertschöpfungsketten
- P-102 Otto Spaniol (Eds.): Mobile Services and Personalized Environments
- P-103 Alfons Kemper, Harald Schöning, Thomas Rose, Matthias Jarke, Thomas Seidl, Christoph Quix, Christoph Brochhaus (Hrsg.): Datenbanksysteme in Business, Technologie und Web (BTW 2007)
- P-104 Birgitta König-Ries, Franz Lehner, Rainer Malaka, Can Türker (Hrsg.): MMS 2007: Mobilität und mobile Informationssysteme
- P-105 Wolf-Gideon Bleek, Jörg Raasch, Heinz Züllighoven (Hrsg.): Software Engineering 2007
- P-106 Wolf-Gideon Bleek, Henning Schwentner, Heinz Züllighoven (Hrsg.): Software Engineering 2007 – Beiträge zu den Workshops
- P-107 Heinrich C. Mayr, Dimitris Karagiannis (eds.) Information Systems Technology and its Applications
- P-108 Arslan Brömme, Christoph Busch, Detlef Hühnlein (eds.) BIOSIG 2007: Biometrics and Electronic Signatures
- P-109 Rainer Koschke, Otthein Herzog, Karl-Heinz Rödiger, Marc Ronthaler (Hrsg.): INFORMATIK 2007 Informatik trifft Logistik Band 1
- P-110 Rainer Koschke, Otthein Herzog, Karl-Heinz Rödiger, Marc Ronthaler (Hrsg.): INFORMATIK 2007 Informatik trifft Logistik Band 2
- P-111 Christian Eibl, Johannes Magenheimer, Sigrid Schubert, Martin Wessner (Hrsg.): DeLFI 2007: 5. e-Learning Fachtagung Informatik
- P-112 Sigrid Schubert (Hrsg.) Didaktik der Informatik in Theorie und Praxis
- P-113 Sören Auer, Christian Bizer, Claudia Müller, Anna V. Zhdanova (Eds.) The Social Semantic Web 2007 Proceedings of the 1st Conference on Social Semantic Web (CSSW)
- P-114 Sandra Frings, Oliver Göbel, Detlef Günther, Hardo G. Hase, Jens Nedon, Dirk Schadt, Arslan Brömme (Eds.) IMF2007 IT-incident management & IT-forensics Proceedings of the 3rd International Conference on IT-Incident Management & IT-Forensics
- P-115 Claudia Falter, Alexander Schliep, Joachim Selbig, Martin Vingron and Dirk Walther (Eds.) German conference on bioinformatics GCB 2007
- P-116 Witold Abramowicz, Leszek Maciszek (Eds.) Business Process and Services Computing 1st International Working Conference on Business Process and Services Computing BPSC 2007
- P-117 Ryszard Kowalczyk (Ed.) Grid service engineering and management The 4th International Conference on Grid Service Engineering and Management GSEM 2007
- P-118 Andreas Hein, Wilfried Thoben, Hans-Jürgen Appelrath, Peter Jensch (Eds.) European Conference on ehealth 2007
- P-119 Manfred Reichert, Stefan Strecker, Klaus Turowski (Eds.) Enterprise Modelling and Information Systems Architectures Concepts and Applications
- P-120 Adam Pawlak, Kurt Sandkuhl, Wojciech Cholewa, Leandro Soares Indrusiak (Eds.) Coordination of Collaborative Engineering - State of the Art and Future Challenges
- P-121 Korbinian Herrmann, Bernd Bruegge (Hrsg.) Software Engineering 2008 Fachtagung des GI-Fachbereichs Softwaretechnik
- P-122 Walid Maalej, Bernd Bruegge (Hrsg.) Software Engineering 2008 - Workshopband Fachtagung des GI-Fachbereichs Softwaretechnik

- P-123 Michael H. Breitner, Martin Breunig, Elgar Fleisch, Ley Pousttchi, Klaus Turowski (Hrsg.)
Mobile und Ubiquitäre Informationssysteme – Technologien, Prozesse, Marktfähigkeit
Proceedings zur 3. Konferenz Mobile und Ubiquitäre Informationssysteme (MMS 2008)
- P-124 Wolfgang E. Nagel, Rolf Hoffmann, Andreas Koch (Eds.)
9th Workshop on Parallel Systems and Algorithms (PASA)
Workshop of the GI/ITG Special Interest Groups PARS and PARVA
- P-125 Rolf A.E. Müller, Hans-H. Sundermeier, Ludwig Theuvsen, Stephanie Schütze, Marlies Morgenstern (Hrsg.)
Unternehmens-IT:
Führungsinstrument oder Verwaltungsbürde
Referate der 28. GIL Jahrestagung
- P-126 Rainer Gimnich, Uwe Kaiser, Jochen Quante, Andreas Winter (Hrsg.)
10th Workshop Software Reengineering (WSR 2008)
- P-127 Thomas Kühne, Wolfgang Reisig, Friedrich Steimann (Hrsg.)
Modellierung 2008
- P-128 Ammar Alkassar, Jörg Siekmann (Hrsg.)
Sicherheit 2008
Sicherheit, Schutz und Zuverlässigkeit
Beiträge der 4. Jahrestagung des Fachbereichs Sicherheit der Gesellschaft für Informatik e.V. (GI)
2.-4. April 2008
Saarbrücken, Germany
- P-129 Wolfgang Hesse, Andreas Oberweis (Eds.)
Sigsand-Europe 2008
Proceedings of the Third AIS SIGSAND European Symposium on Analysis, Design, Use and Societal Impact of Information Systems
- P-130 Paul Müller, Bernhard Neumair, Gabi Dreö Rodosek (Hrsg.)
1. DFN-Forum Kommunikationstechnologien Beiträge der Fachtagung
- P-131 Robert Krimmer, Rüdiger Grimm (Eds.)
3rd International Conference on Electronic Voting 2008
Co-organized by Council of Europe, Gesellschaft für Informatik and E-Voting.CC
- P-132 Silke Seehusen, Ulrike Lucke, Stefan Fischer (Hrsg.)
DeLFI 2008:
Die 6. e-Learning Fachtagung Informatik
- P-133 Heinz-Gerd Hegering, Axel Lehmann, Hans Jürgen Ohlbach, Christian Scheideler (Hrsg.)
INFORMATIK 2008
Beherrschbare Systeme – dank Informatik Band 1
- P-134 Heinz-Gerd Hegering, Axel Lehmann, Hans Jürgen Ohlbach, Christian Scheideler (Hrsg.)
INFORMATIK 2008
Beherrschbare Systeme – dank Informatik Band 2
- P-135 Torsten Brinda, Michael Fothe, Peter Hubwieser, Kirsten Schlüter (Hrsg.)
Didaktik der Informatik – Aktuelle Forschungsergebnisse
- P-136 Andreas Beyer, Michael Schroeder (Eds.)
German Conference on Bioinformatics GCB 2008
- P-137 Arslan Brömme, Christoph Busch, Detlef Hühnlein (Eds.)
BIOSIG 2008: Biometrics and Electronic Signatures
- P-138 Barbara Dinter, Robert Winter, Peter Chamoni, Norbert Gronau, Klaus Turowski (Hrsg.)
Synergien durch Integration und Informationslogistik
Proceedings zur DW2008
- P-139 Georg Herzwurm, Martin Mikusz (Hrsg.)
Industrialisierung des Software-Managements
Fachtagung des GI-Fachausschusses Management der Anwendungsentwicklung und -wartung im Fachbereich Wirtschaftsinformatik
- P-140 Oliver Göbel, Sandra Frings, Detlef Günther, Jens Nedon, Dirk Schadt (Eds.)
IMF 2008 - IT Incident Management & IT Forensics
- P-141 Peter Loos, Markus Nüttgens, Klaus Turowski, Dirk Werth (Hrsg.)
Modellierung betrieblicher Informationssysteme (MobIS 2008)
Modellierung zwischen SOA und Compliance Management
- P-142 R. Bill, P. Korduan, L. Theuvsen, M. Morgenstern (Hrsg.)
Anforderungen an die Agrarinformatik durch Globalisierung und Klimaveränderung
- P-143 Peter Liggesmeyer, Gregor Engels, Jürgen Münch, Jörg Dörr, Norman Riegel (Hrsg.)
Software Engineering 2009
Fachtagung des GI-Fachbereichs Softwaretechnik

- P-144 Johann-Christoph Freytag, Thomas Ruf, Wolfgang Lehner, Gottfried Vossen (Hrsg.)
Datenbanksysteme in Business, Technologie und Web (BTW)
- P-145 Knut Hinkelmann, Holger Wache (Eds.)
WM2009: 5th Conference on Professional Knowledge Management
- P-146 Markus Bick, Martin Breunig, Hagen Höpfner (Hrsg.)
Mobile und Ubiquitäre Informationssysteme – Entwicklung, Implementierung und Anwendung
4. Konferenz Mobile und Ubiquitäre Informationssysteme (MMS 2009)
- P-147 Witold Abramowicz, Leszek Maciaszek, Ryszard Kowalczyk, Andreas Speck (Eds.)
Business Process, Services Computing and Intelligent Service Management
BPSC 2009 · ISM 2009 · YRW-MBP 2009
- P-148 Christian Erfurth, Gerald Eichler, Volkmar Schau (Eds.)
9th International Conference on Innovative Internet Community Systems
I²CS 2009
- P-149 Paul Müller, Bernhard Neumair, Gabi Dreö Rodosek (Hrsg.)
2. DFN-Forum
Kommunikationstechnologien
Beiträge der Fachtagung
- P-150 Jürgen Münch, Peter Liggesmeyer (Hrsg.)
Software Engineering
2009 - Workshopband
- P-151 Armin Heinzl, Peter Dadam, Stefan Kirn, Peter Lockemann (Eds.)
PRIMIUM
Process Innovation for Enterprise Software
- P-152 Jan Mendling, Stefanie Rinderle-Ma, Werner Esswein (Eds.)
Enterprise Modelling and Information Systems Architectures
Proceedings of the 3rd Int'l Workshop EMISA 2009
- P-153 Andreas Schwill, Nicolas Apostolopoulos (Hrsg.)
Lernen im Digitalen Zeitalter
DeLFI 2009 – Die 7. E-Learning Fachtagung Informatik
- P-154 Stefan Fischer, Erik Maehle, Rüdiger Reischuk (Hrsg.)
INFORMATIK 2009
Im Focus das Leben
- P-155 Arslan Brömme, Christoph Busch, Detlef Hühnlein (Eds.)
BIOSIG 2009:
Biometrics and Electronic Signatures
Proceedings of the Special Interest Group on Biometrics and Electronic Signatures
- P-156 Bernhard Koerber (Hrsg.)
Zukunft braucht Herkunft
25 Jahre »INFOS – Informatik und Schule«
- P-157 Ivo Grosse, Steffen Neumann, Stefan Posch, Falk Schreiber, Peter Stadler (Eds.)
German Conference on Bioinformatics
2009
- P-158 W. Claupein, L. Theuvsen, A. Kämpf, M. Morgenstern (Hrsg.)
Precision Agriculture
Reloaded – Informationsgestützte Landwirtschaft
- P-159 Gregor Engels, Markus Luckey, Wilhelm Schäfer (Hrsg.)
Software Engineering 2010
- P-160 Gregor Engels, Markus Luckey, Alexander Pretschner, Ralf Reussner (Hrsg.)
Software Engineering 2010 – Workshopband
(inkl. Doktorandensymposium)
- P-161 Gregor Engels, Dimitris Karagiannis, Heinrich C. Mayr (Hrsg.)
Modellierung 2010
- P-162 Maria A. Wimmer, Uwe Brinkhoff, Siegfried Kaiser, Dagmar Lück-Schneider, Erich Schweighofer, Andreas Wiebe (Hrsg.)
Vernetzte IT für einen effektiven Staat
Gemeinsame Fachtagung
Verwaltungsinformatik (FTVI) und
Fachtagung Rechtsinformatik (FTRI) 2010
- P-163 Markus Bick, Stefan Eulgem, Elgar Fleisch, J. Felix Hampe, Birgitta König-Ries, Franz Lehner, Key Pousttchi, Kai Rannenberg (Hrsg.)
Mobile und Ubiquitäre Informationssysteme
Technologien, Anwendungen und Dienste zur Unterstützung von mobiler Kollaboration
- P-164 Arslan Brömme, Christoph Busch (Eds.)
BIOSIG 2010: Biometrics and Electronic Signatures
Proceedings of the Special Interest Group on Biometrics and Electronic Signatures

- P-165 Gerald Eichler, Peter Kropf,
Ulrike Lechner, Phayung Meesad,
Herwig Unger (Eds.)
10th International Conference on
Innovative Internet Community Systems
(I²CS) – Jubilee Edition 2010 –
- P-166 Paul Müller, Bernhard Neumair,
Gabi Dreö Rodosek (Hrsg.)
3. DFN-Forum Kommunikationstechnologien
Beiträge der Fachtagung
- P-167 Robert Krimmer, Rüdiger Grimm (Eds.)
4th International Conference on
Electronic Voting 2010
co-organized by the Council of Europe,
Gesellschaft für Informatik and
E-Voting.CC
- P-168 Ira Diethelm, Christina Dörge,
Claudia Hildebrandt,
Carsten Schulte (Hrsg.)
Didaktik der Informatik
Möglichkeiten empirischer
Forschungsmethoden und Perspektiven
der Fachdidaktik
- P-169 Michael Kerres, Nadine Ojstersek
Ulrik Schroeder, Ulrich Hoppe (Hrsg.)
DeLFI 2010 - 8. Tagung
der Fachgruppe E-Learning
der Gesellschaft für Informatik e.V.
- P-170 Felix C. Freiling (Hrsg.)
Sicherheit 2010
Sicherheit, Schutz und Zuverlässigkeit
- P-171 Werner Esswein, Klaus Turowski,
Martin Jührisch (Hrsg.)
Modellierung betrieblicher
Informationssysteme (MobIS 2010)
Modellgestütztes Management
- P-172 Stefan Klink, Agnes Koschmider
Marco Mevius, Andreas Oberweis (Hrsg.)
EMISA 2010
Einflussfaktoren auf die Entwicklung
flexibler, integrierter Informationssysteme
Beiträge des Workshops der GI-
Fachgruppe EMISA
(Entwicklungsmethoden für Infor-
mationssysteme und deren Anwendung)
- P-173 Dietmar Schomburg,
Andreas Grote (Eds.)
German Conference on Bioinformatics
2010
- P-174 Arslan Brömme, Torsten Eymann,
Detlef Hühnlein, Heiko Roßnagel,
Paul Schmücker (Hrsg.)
perspeGKtive 2010
Workshop „Innovative und sichere
Informationstechnologie für das
Gesundheitswesen von morgen“
- P-175 Klaus-Peter Fährnich,
Bogdan Franczyk (Hrsg.)
INFORMATIK 2010
Service Science – Neue Perspektiven für
die Informatik
Band 1
- P-176 Klaus-Peter Fährnich,
Bogdan Franczyk (Hrsg.)
INFORMATIK 2010
Service Science – Neue Perspektiven für
die Informatik
Band 2
- P-177 Witold Abramowicz, Rainer Alt,
Klaus-Peter Fährnich, Bogdan Franczyk,
Leszek A. Maciaszek (Eds.)
INFORMATIK 2010
Business Process and Service Science –
Proceedings of ISSS and BPSC
- P-178 Wolfram Pietsch, Benedikt Krams (Hrsg.)
Vom Projekt zum Produkt
Fachtagung des GI-Fachausschusses
Management der
Anwendungsentwicklung und -wartung
im Fachbereich Wirtschaftsinformatik
(WI-MAW), Aachen, 2010
- P-179 Stefan Gruner, Bernhard Rumpe (Eds.)
FM+AM 2010
Second International Workshop on Formal
Methods and Agile Methods
- P-180 Theo Härder, Wolfgang Lehner,
Bernhard Mitschang, Harald Schöning,
Holger Schwarz (Hrsg.)
Datenbanksysteme für Business,
Technologie und Web (BTW)
14. Fachtagung des GI-Fachbereichs
„Datenbanken und Informationssysteme“
(DBIS)
- P-181 Michael Clasen, Otto Schätzel,
Brigitte Theuvsen (Hrsg.)
Qualität und Effizienz durch
informationsgestützte Landwirtschaft,
Fokus: Moderne Weinwirtschaft
- P-182 Ronald Maier (Hrsg.)
6th Conference on Professional
Knowledge Management
From Knowledge to Action
- P-183 Ralf Reussner, Matthias Grund, Andreas
Oberweis, Walter Tichy (Hrsg.)
Software Engineering 2011
Fachtagung des GI-Fachbereichs
Softwaretechnik
- P-184 Ralf Reussner, Alexander Pretschner,
Stefan Jähnichen (Hrsg.)
Software Engineering 2011
Workshopband
(inkl. Doktorandensymposium)

- P-185 Hagen Höpfner, Günther Specht,
Thomas Ritz, Christian Bunse (Hrsg.)
MMS 2011: Mobile und ubiquitäre
Informationssysteme Proceedings zur
6. Konferenz Mobile und Ubiquitäre
Informationssysteme (MMS 2011)
- P-187 Paul Müller, Bernhard Neumair,
Gabi Dreo Rodosek (Hrsg.)
4. DFN-Forum Kommunikations-
technologien, Beiträge der Fachtagung
20. Juni bis 21. Juni 2011 Bonn

The titles can be purchased at:

Köllen Druck + Verlag GmbH

Ernst-Robert-Curtius-Str. 14 · D-53117 Bonn

Fax: +49 (0)228/9898222

E-Mail: druckverlag@koellen.de