

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 978-3-88579-625-1

The 7. DFN-Forum Communication Technologies 2014 is taking place in Fulda, Germany, from June 16<sup>th</sup> to June 17<sup>th</sup>.

This volume contains 13 papers selected for presentation at the conference.

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



P. Müller, B. Neumair, H. Reiser, G. Dreö Rodosek (Hrsg.): 7. DFN-Forum 2014

# GI-Edition

## Lecture Notes in Informatics

**Paul Müller, Bernhard Neumair,  
Helmut Reiser, Gabi Dreö Rodosek (Hrsg.)**

## 7. DFN-Forum Kommunikations- technologien

**16.–17. Juni 2014  
Fulda**







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

## **7. DFN-Forum Kommunikationstechnologien**

**Beiträge der Fachtagung**

**16.06. – 17.06.2014  
Fulda**

Gesellschaft für Informatik e.V. (GI)

## **Lecture Notes in Informatics (LNI) - Proceedings**

Series of the Gesellschaft für Informatik (GI)

Volume P-231

ISBN 978-3-88579-625-1

ISSN 1617-5468

### **Volume Editors**

Prof. Dr. Paul Müller (pmueller@informatik.uni-kl.de)

Technische Universität Kaiserslautern

Postfach 3049, 67653 Kaiserslautern

Prof. Dr. Bernhard Neumair (bernhard.neumair@kit.edu)

Karlsruher Institut für Technologie (KIT)

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

PD Dr. Helmut Reiser (reiser@lrz.de)

Leibniz-Rechenzentrum

Boltzmannstraße 1, 85748 Garching

Prof. Dr. Gabi Dreö Rodosek (Gabi.Dreö@unibw.de)

Universität der Bundeswehr München

Werner-Heisenberg-Weg 39, 85577 Neubiberg

### **Series Editorial Board**

Heinrich C. Mayr, Alpen-Adria-Universität Klagenfurt, Austria

(Chairman, mayr@ifit.uni-klu.ac.at)

Dieter Fellner, Technische Universität Darmstadt, Germany

Ulrich Flegel, Hochschule für Technik, Stuttgart, Germany

Ulrich Frank, Universität Duisburg-Essen, Germany

Johann-Christoph Freytag, Humboldt-Universität zu Berlin, 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

Peter Sanders, Karlsruher Institut für Technologie (KIT), Germany

Sigrid Schubert, Universität Siegen, Germany

Ingo Timm, Universität Trier, Germany

Karin Vosseberg, Hochschule Bremerhaven, Germany

Maria Wimmer, Universität Koblenz-Landau, 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 2014

**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 7. DFN-Forum Kommunikationstechnologien „Verteilte Systeme im Wissenschaftsbereich“ steht in dieser Tradition. Nach den sehr erfolgreichen Vorgängerveranstaltungen in Kaiserslautern, München, Konstanz, Bonn, Regensburg und Erlangen wird die diesjährige Tagung vom DFN-Verein und der Hochschule Fulda gemeinsam mit dem ZKI e.V. und der GI am 16. und 17. Juni 2014 in Fulda 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 sich neben Fragen der Energieeffizienz und der Green IT intensiv mit Virtualisierung und Cloud Services befasst. Traditionell ist auch die IT-Sicherheit ein wichtiges Thema des Forums, das in jüngster Zeit durch die NSA-Affäre enorme Aufmerksamkeit bekommen hat. Bei einer Podiumsdiskussion zum Thema „Mehr Sicherheit durch Schengen-Routing und Deutsches Internet: politischer Aktionismus oder realisierbare Vision?“ werden Experten die durch die Affäre initiierten Aktivitäten kritisch beleuchten und einen Blick in die Zukunft des „deutschen Internets“ eröffnen. Ergänzt wird dies durch einen eingeladenen Beitrag zum Internet Exchange in Frankfurt im Lichte immer größerer Bandbreiten. Abgerundet wird das Programm durch mehrere Beiträge zum Management von Forschungsdaten und einen eingela-

denen Vortrag zum EU-Projekt EUDAT und Diensten für eine kollaborative Dateninfrastruktur.

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

Fulda, April 2014

Gabi Dreö Rodosek  
Paul Müller  
Bernhard Neumair  
Helmut Reiser

## **Programmkomitee**

*Rainer Bockholt*, Universität Bonn

*Alexander Clemm*, Cisco USA

*Gabriele Dobler*, Universität Erlangen-Nürnberg

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

*Thomas Eickermann*, Forschungszentrum Jülich

*Alfred Geiger*, T-Systems Sfr

*Wolfgang Gentzsch*, The UberCloud

*Hannes Hartenstein*, KIT

*Peter Holleczech*, Universität Erlangen-Nürnberg

*Eike Jessen*, Technische Universität München

*Peter Klingebiel*, Hochschule Fulda

*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

*Helmut Reiser (Co-Chair)*, LRZ München

*Sebastian Rieger*, Hochschule Fulda

*Harald Roelle*, Siemens AG

*Peter Schirnbacher*, Humboldt-Universität zu Berlin

*Uwe Schwiegelshohn*, TU Dortmund

*Manfred Seedig*, Universität Kassel

*Marcel Waldvogel*, Universität Konstanz

*René Wies*, BMW Group

*Martin Wimmer*, Universität Regensburg





# Inhaltsverzeichnis

## Management von Forschungsdaten

|                                                                                                                                                                                                               |    |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Stefan Buddenbohm, Harry Enke, Matthias Hofmann, Jochen Klar, Heike Neuroth, Uwe Schwiegelshohn</b><br><i>A Life Cycle Model für Collaborative Research Environments</i> .....                             | 1  |
| <b>Jan Potthoff, Jos van Wezel, Matthias Razum, Marius Walk</b><br><i>Anforderungen eines nachhaltigen, disziplinübergreifenden Forschungsdaten-Repositorys</i> .....                                         | 11 |
| <b>Johannes Fricke, Andreas Grandel</b><br><i>Langzeitarchivierung von Forschungsdaten im Sonderforschungsgebiet 840 „von partikulären Nanosystemen zur Mesotechnologie“ an der Universität Bayreuth</i> .... | 21 |

## Energieeffizienz, Green-IT

|                                                                                                                                                          |    |
|----------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Gudrun Oevel, Sebastian Porombka, Maximilian Boehner</b><br><i>Energieeffizienz im WLAN</i> .....                                                     | 33 |
| <b>Kai Spindler, Sebastian Rieger</b><br><i>AEQUO – Adaptive und energieeffiziente Verteilung von virtuellen Maschinen in OpenStack-Umgebungen</i> ..... | 45 |

## Virtualisierung und Cloud Services

|                                                                                                                                                     |    |
|-----------------------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Felix Maurer, Sebastian Labitzke</b><br><i>FOSP: Towards a Federated Object Sharing Protocol that Unifies Operations on Social Content</i> ..... | 57 |
| <b>Christopher Ritter, Patrick Bittner, Odej Kao</b><br><i>Vom BYOE zu GYSE</i> .....                                                               | 67 |
| <b>Konrad Meier, Steffen Philipp</b><br><i>Anonymisierung und externe Speicherung in Cloud-Speichersystemen</i> .....                               | 75 |
| <b>Martin Metzker, Dieter Kranzlmüller</b><br><i>Mapping virtual paths in virtualization environments</i> .....                                     | 87 |

IT-Sicherheit

**Mario Golling, Robert Koch, Peter Hillmann, Gabi Dreo Rodosek**  
*Smart Defence: An Architecture for new Challenges to Cyber Security* ..... 99

**Felix Eye, Wolfgang Hommel, Helmut Reiser**  
*Herausforderungen und Anforderungen für ein organisationsweites Notfall-Passwortmanagement*..... 109

**Marcel Waldvogel, Jürgen Kolle**  
*SIEGE: Service-independent enterprise-grade protection against password scans* ..... 121

**Mario Golling, Robert Koch, Lars Stiemert**  
*Architektur zur mehrstufigen Angriffserkennung in Hochgeschwindigkeits-Backbone-Netzen*..... 131

# **Management von Forschungsdaten**



# A Life Cycle Model for Collaborative Research Environments

Stefan Buddenbohm,<sup>†</sup> Harry Enke,<sup>‡</sup> Matthias Hofmann\*  
Jochen Klar,<sup>‡</sup> Heike Neuroth,<sup>†</sup> Uwe Schwiegelshohn\*

**Abstract:** Virtual or Collaborative Research Environments are important pillars of modern research in many scientific disciplines. Therefore, the development and particularly the sustained operation of these environments has become an important challenge in the area of modern research infrastructures. In order to support communities in dealing with this challenge including design and evaluation, we have developed a life cycle model for Collaborative Research Environments. Based on this life cycle model, we identify common pitfalls and suggest a generic catalog of criteria to be used by developers, operators, and funding agencies. To consider the heterogeneity of Collaborative Research Environments, additional discipline specific criteria can be incorporated.

## 1 Introduction

Across many disciplines modern science often requires large amounts of digital research data and various software products. Therefore, information technology has become vital for most research environments. Contrary to expectations published in many magazines the use of IT in research has led to a significant increase of expenses during the last years, which are not compensated by technological progress. But IT also offers new opportunities as it enables us to access the components of our research environments from most locations of the world due to the availability of high bandwidth connectivity almost everywhere. This location-independent access allows us to share the components of our research environment with many colleagues throughout the world and thus leads to new and more efficient research collaborations. We can also distribute the costs of maintaining large data centres or specialized research software across many different interest groups. These synergy effects offer the potential of at least partially compensating the rising costs of science. As such research environments are shared by research groups belonging to different research institutions, it is often called a virtual research environment to emphasize that the cooperation of people from different institutions forms a *virtual organization*. Unfortunately, this this name is misleading as the word ‘virtual’ is often regarded as the opposite of ‘real’. But any virtual research environment is real for both the user and the provider of

---

\* {matthias.hofmann, uwe.schwiegelshohn}@udo.edu, Robotics Research Institute, TU Dortmund University

† {buddenbohm, neuroth}@sub.uni-goettingen.de, Göttingen State and University Library

‡ {henke, jklar}@aip.de, Leibniz Institute for Astrophysics Potsdam (AIP)

the environment. Since the mentioned connectivity also supports communication between different research groups working on the same topic, we prefer the expression *collaborative research environment (CRE)* although collaboration does not comprise all uses of the research environment.

As the potential benefits of collaborative research environments are recognized in many disciplines, a significant number of such research environments has been created recently. But we must realize that despite the interest in collaborative research environments and the growing number of new collaborative research environments, sustained operation of a collaborative research environment is rarely achieved. In a project supported by the German research council, we have examined collaborative research environments in various research disciplines and related IT infrastructures. This study has led us to the characterization of different phases in the life cycle of a collaborative research environment. We are able to specify different pitfalls in these phases. There are various reasons for these obstacles: sometimes the effort to transfer a prototype environment into a production environment is underestimated, sometimes legal and administrative procedures hamper establishing an environment that is shared by different institutions in different states or countries. We strongly believe that only a thorough understanding of the complete life cycle of a research environment will help to overcome these obstacles. Therefore, we will explain this life cycle and the different aspects of its phases and connected success criteria in this paper.

The paper is organized as follows. While Section 2 deals with related work, Section 3 describes the methodology that has been applied. In Section 4, we suggest a life cycle model for CREs, and describe its characteristics. The paper concludes in Section 5.

## 2 Related Work

Our work follows up on numerous studies which have been recently performed in this field, whereby we focus for obvious reasons on the discussion in Germany. First and foremost the work group "Virtual Research Environments" of the Priority Initiative Digital Information by the Alliance of German Science Organisations[dAIDI11] is of importance since we use their discussions and studies as preliminary work for our approach, especially their concept to assess the degree of maturity or establishment of CREs within their respective communities. Based on the work of the 'Joint Information Systems Committee' (JISC) in the course of the 'British Virtual Research Environment Programme'[CR10], it became clear to us that we must choose a pragmatic approach emphasizing the demands of researchers and developers/ operators.

An alternative approach could have focused on the funders' perspective and the evaluative aspects of measuring research infrastructures, e.g. to facilitate funding decisions, like the ERINA, and ERINA+<sup>1</sup>-projects funded by the European Commission within the FP7-Programme. Decisive criteria for ERINA+ were the satisfaction of the user community and the amount of CRE-assignable results.

---

<sup>1</sup><http://www.erinaplus.eu/>

The lack of comparability of the results remains an unsolved problem for all these approaches. Although it is possible to evaluate the quality of CREs based on reproducible facts, it is in our eyes impossible (or at least questionable) to compare these evaluations.

Furthermore, the work group ‘Virtual Research Environments’ of the German Initiative for Network Information (DINI)[For11], the DFG project ‘Radieschen’[For13b], and the WissGrid project of the German D-Grid Initiative[LE13] were important for our work as they contributed valuable experiences from various past and current projects.

Our modular CRE life cycle entailed fruitful discussions in various interviews and workshops (see Section 3). Although some generic aspects may also fit into software engineering models, CRE specific inspiration was specifically derived from the work of the ‘Co-operative University Computing Facilities’ (SURF) with its VRE Starters Kits[SUR12].

### 3 Methodology

Ideally a collaborative research environment supports the complete research process or at least major parts of it. Since CREs are relevant to many different scientific communities and there are significant differences in the research process between different communities including tools and methodological approaches, we cannot expect that a single scheme for a CRE life cycle model and success criteria fits all communities. But due to the efforts in time and money required to establish a successful CRE, we also cannot afford to build every CRE from scratch. Instead, we must identify common properties and processes within CREs to allow guidance and evaluation how to generate and operate a CRE. This is in particular of interest for operators of CRE and funding agencies.

Similar to large software systems there are different approaches to build and run a CRE. Instead of using a theoretical approach, we decided on collecting information from existing CREs or related systems by conducting expert interviews with suitable stakeholders based on our life cycle model (LCM). In addition to institutions that operate CREs, also e-Infrastructures and software providers such as European Grid Infrastructure (EGI), the German National Research and Education Network (DFN), and the Higher Education Information System (HIS) provided valuable input as they face problems similar to CRE, especially regarding financial issues. The interviews were conducted over a period of three months. Since our interview partners provided expertise in different areas, we tailored the interviews individually instead of using a fixed scheme. As some interviews have yielded results that were surprising to us, we decided to incorporate such results into the questions of later interviews to determine whether these results are important in general or are only relevant for a particular CRE.

Additional input was collected within two workshops. The first workshop reviewed the results of the interviews. To this end, all interview partners and additional representatives from different scientific communities, CRE operators, and funding agencies were invited. The second workshop was organized by the German Research Council and discussed organizational and management issues for virtual research environments.

As a result of our study, we have developed a life cycle model that is suitable for most



CREs. This LCM comprises five phases (see Section 4) and was considered to be very useful by the participants of the first workshop where it was presented for the first time. We structure the main part of this paper according to the phases of the model. However, it is noteworthy that the establishment of a CRE might not strictly follow the very life cycle, i.e. iterations between individual phases like in current software development models occur. Thus, the life cycle model describes an idealized view that aims to guide through the establishment and operation of a CRE.

Regarding the evaluation of a CRE or - to be less controversial - to determine best practices for their generation and operation of a CRE, we propose to define a generic set of core criteria and milestones. Key objective is the usefulness of the CRE to researchers and its efficiency in operation. Due to the heterogeneity of CREs, we suggest that the generic success criteria serve as a pool from which an individual set must be selected in each case depending on the scope and service portfolio of the specific CRE. Furthermore, it is necessary to supplement these generic criteria with discipline-specific or community aspects and criteria since, for instance, a CRE in the arts and humanities at least partially requires different criteria than a CRE in engineering or natural sciences.

## 4 Life Cycle of a Collaborative Research Environment

The life of a CRE can be divided into five different phases. The schematic sequence is displayed in Fig. 1 while a list of success criteria is given in Table 1. In the first phase, the *prototype*, a first version of the CRE is constructed and deployed for a limited number of users. This phase can demonstrate the overall feasibility of the system and provide first assessment of the acceptance by the corresponding scientific community. If the prototype is, as we will discuss later, considered successful, the subsequent phase of *development* is used to address all aspects, which are crucial for a sustainable version of the CRE, but were left out while developing the prototype. An operational concept is most important in this context, the services of the CRE are matched with the required resources and a budget in which the financial resources and the corresponding funding sources are specified.

Again, if a set of criteria and milestones is met, the CRE enters the next phase. For a more infrastructure-like CRE we call this phase *operation*, while for exclusive software projects the term *product* is better suited. In this phase the CRE is made available to the corresponding scientific community. During this time, a special focus lies on the support of the users of the CRE and the subsequent technical and organizational improvement. When, after some time of operation, the decision to discontinue the CRE is made, a phase of *transfer* is used to move the components of the CRE which are worth conserving (e.g. research data, source code, etc.) to other (infra-)structures. The remaining technical and organizational components of the CRE are then dissolved in the phase of *liquidation*.

In general, there are two different ways of operating a CRE. Either the CRE is built and operated by a collaboration as a separate infrastructure unit. This case is mainly considered in the following. Or the CRE is embedded as a module in a broader infrastructural context, e.g. as a part of a data center. Here, prototype and development phases, as well as transfer

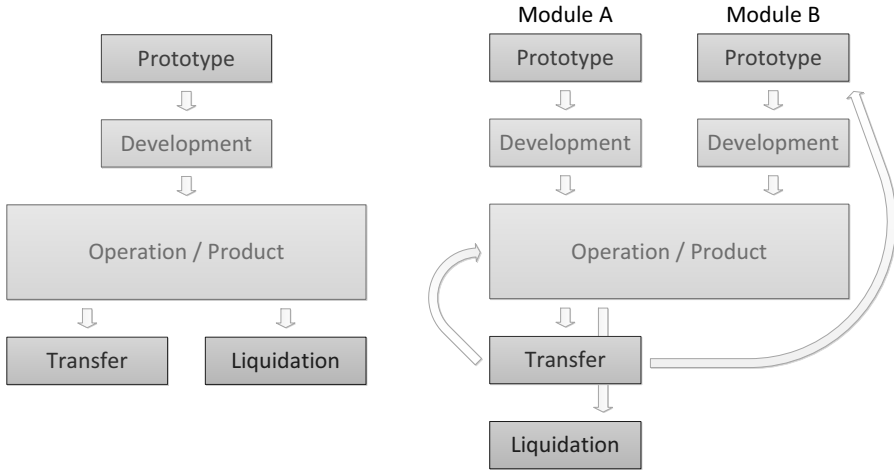


Figure 1: The life cycle of a monolithic CRE is shown in the left figure. After the creation of a prototype and a subsequent development phase, the phase of operation (infrastructure) or product (software) is reached. When the CRE is discontinued, components of the CRE which are worth retaining are transferred, while the rest is liquidated. The right figure shows the complex life cycle of a modular CRE. A first module A is created similar to the first scenario and passes the prototype and the development phases. Then it enters a operational phase in a shared infrastructure with other modules. At the end of its lifetime some of its components are transferred to modules still in the phase of operation, while others are used in the creation of a new module B. The remain of module A is dissolved in the phase of liquidation.

and liquidation follow the same rules, but are applied to this specific module only. As shown in Fig. 1, parts of a CRE can also be transferred not only to external infrastructures but may be reused in all stages of a new CRE module.

#### 4.1 The Prototype

The first step towards the implementation of a new CRE is the creation of a prototype. In general this is carried out in form of a research project including the usual review process. Accordingly, one of the first steps (usually part of the initial application) is to specify the internal organizational structure of the project and how the work is distributed. The first steps of the project itself then comprises a set of *preliminary studies* concerning the environment of the new CRE. Those studies include a detailed *demand analysis* as well as a *determination of the relation* of the new CRE within the context of existing infrastructures, both locally (e.g. computing centres, libraries) and with respect to the main research community of the CRE (e.g. data centres). Moreover, a *positioning* with respect to the other stakeholders in the research community of the CRE, a *thorough analysis* of scientific work flows which the CRE aims to support, and a *market analysis* of the considered hardware and software components are needed.

We consider these preliminary studies to be success criteria in a broader sense; they are essential for progressing in the further development of the CRE. Building upon these analyses, the further process is split into two branches, a technical branch in which the technical infrastructure is set up and the software development takes place, and an organizational branch, which is responsible for the development of corresponding organizational structures. In our model, these branches are further divided into specification and implementation.

During the *technical specification*, the services of the CRE are defined based on the preceding analysis of the scientific work flows. The software is being developed using a well-known, established process which is usually the waterfall model. Moreover, standards and as metadata are chosen, and decisions for the used hardware and software are made. In the subsequent *technical implementation*, the software development work is carried out, and the installation of hardware and software takes place. Extensive testing of the system (i.e. unit, integration tests) ensure a running system.

The *organizational specification* comprises a round-up of the different access permissions for the resources of the CRE (e.g. different layers of access to research data) as well as possible policies. In addition, strategies for integration to the community and the creation of support facilities such as a help desk and training group need to be prepared. The *organizational implementation* phase is then used to raise awareness in the corresponding research field and build a community around the new CRE. The set-up of preliminary support facilities and first training sessions for future power-user can already take place at this point.

The described works culminate in a *pilot system* in which the new prototype is made available to a limited number of scientific users. This is not only used for further testing and refinement of the system, but also offers the possibility to adjust the CRE through smaller improvements to meet the needs of its community.

In order to gain information on the particular technical, financial and organizational requirements, especially with respect to a sustainable modus of operation, all described phases are accompanied by a systematic monitoring, which is part of the project management.

## 4.2 Development Phase

During the pilot phase of the prototype, first impressions of the working system and in particular an assessment of the user acceptance can be done. Based on this work, a decision for the proceeding of the next phase of the CRE, the development phase, is made.

With respect to financing CREs, it is beneficial to distinguish between the first two phases (i.e. *prototype* and *development phase*) that build up or extend the very product (CRE), and the operation phase. While the build up phases are often financially supported in form of projects, the operation phase requires a sophisticated business model to ensure potentially long-term provisioning of services. Therefore, it is of vital importance to integrate a business plan into the planning efforts at this early stage.

The phase of development is then used to prepare all necessary features the CRE needs for a subsequent sustainable phase of operation, but which are generally not implemented while creating the prototype. This includes the creation of a *legal entity* for the CRE and a *concept of operations* in which the provided services to the community are related to the required resources. Also a financial budget needs to be prepared. Another important task is the compilation of a sustainability concept, which specifies the whereabouts of different components (e.g. research data, software) of the CRE after the end of its lifetime.

### 4.3 Operation

Besides the technical operation of the CRE, key aspects in this phase are the support of the community and the continuous improvement of the technical and the organizational environment. While no strict temporal sequence of tasks is given, a set of equally important tasks can be identified. An overview of these tasks is given in Fig. 2. In general the critical relevance of user involvement and community acceptance for the success of the collaborative research environment cannot be underestimated. The presented life cycle phases and success criteria consider this aspect by clearly defined tasks (monitoring, periodical evaluation, feedback from users), by preliminary studies before the start of the development phase and later on by dedicated criteria that allow an evaluation of user involvement and community acceptance.

Support for the users is provided by a help desk and a sufficient documentation of the CRE. Through focused events (e.g. workshops, training) the community of the CRE is maintained and expanded. Public relations advertise the CRE in the scientific community and to a wider audience. Although, the CRE should provide a stable working environment for the users, it is necessary to perform further work on the system. Especially requests from users and new features emanating from new user groups require substantial development. Also security issues of the used software components and following updates need to be taken into account. The coordination of these aspects lies within the requirements management. While larger improvements are organized as separate projects that might follow a similar life cycle model as the CRE, smaller improvements and bug fixes can be performed within the structures of the CRE. In addition to the production system, a CRE should have a development and a testing system to integrate developments seamlessly into operation. In conclusion the criteria for the success of a CRE during the operational phase can be summarized as following:

*Community-related aspects*, e.g. the acceptance of the CRE within the potential community of users, the degree to which the CRE meets the specific demands of the users or its ability to improve its services alongside a moderated process of communication between users and developers of the CRE.

*Performance-related aspects*, e.g. the ease of use (not just usability) of the CRE, its preferably seamless integration into already existing research environments of the users and its stability and operational availability.

*Result-related aspects*, e.g. the exploitation of synergies within the research work of the

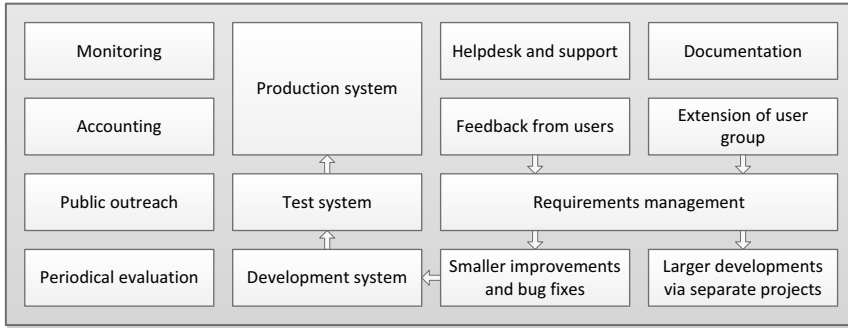


Figure 2: Central aspects of the operation (or product) phase of a CRE.

users. Based on our interviews during the course of the project the synergies are at the forefront of many users' expectations towards a CRE and not the enabling of new research questions with the help of technology. The result-related aspects emphasize the character of the CRE as a research tool for its users, as a means to ease their daily work. Furthermore aspects of collaboration and the use of distributed resources are to be subsumed here.

*Transfer-related aspects*, e.g. the occurrence and use of initial resources, experiences services, and products in other CREs or research environments.

*Competence-related aspects*, e.g. acquisition and transfer of individual and institutional expertise. These can be represented by staff as well as by documented experiences of infrastructure operators and developers that facilitate the development of subsequent services and products. It is also relevant how much a CRE affects the research practice of the respective discipline.

Of course the aforementioned remarks only represent a broad outline of the possible success criteria for a CRE. Each of the five *aspects* must be differentiated into detailed criteria, which are - ideally - as meaningful as possible. For each individual CRE a set of specific criteria must be selected out of the above mentioned aspects. Case-by-case these more generic criteria must be supplemented by discipline-specific criteria depending on the scope and services of the CRE.

Sustainable operation of CREs means that services should run as long as they are required by scientific communities as an integral part of their research. Regarding costs and financing, we clearly identify expenses for scientific and technical staff to be most relevant as a result of the expert interviews.

#### 4.4 Transfer and Liquidation

At one point in time the CRE, or a smaller part of a modular CRE infrastructure, reaches the end of its lifetime, and the decision of its dissolution is made. In this case, the most im-

| Phase               | Criterion                        |
|---------------------|----------------------------------|
| Prototype           | Demand analysis                  |
|                     | Technological landscape study    |
|                     | Technologies / Tools / Standards |
|                     | Monitoring                       |
|                     | Community acceptance             |
| Development phase   | Concept of operations            |
|                     | Organizational form              |
|                     | Legal entity                     |
|                     | Financial budget                 |
| Operation / Product | User acceptance and satisfaction |
|                     | Enabling collaboration           |
|                     | Support / Community              |
|                     | Publications / Results           |
|                     | Build-up of Expertise            |
|                     | Public outreach                  |
| Transfer            | Recurring evaluations            |
|                     | Knowledge transfer               |
|                     | Technology re-use                |

Table 1: Generic success criteria for CREs along the life cycle.

portant task is to transfer any components which are worth preserving to other infrastructures. These infrastructures can be external like community data centres or other CREs, or, in case of a modular CRE, different modules of the same system (see Fig. 1). Components which should be transferred are *research data* (which should be published or migrated to other data centres or CREs), *services* (which should be migrated to other providers), *software* (which should be published open-source or at least handed over to partners who can support them in the future), and *different other deliverables* (which should be published).

A transfer of these components can be mandatory, for instance, due to legal obligations or other commitments like the rules for safeguarding good scientific practice of the German Research Foundation (DFG)[For13a], but there can also be other reasons, like a still existing user community.

During the liquidation, the organizational structures of the CRE are dissolved. If required a successor organization can be founded. A migration support should help the last remaining users with the transition to different infrastructures.

## 5 Conclusion and Outlook

In this paper we described a Life Cycle Model for Collaborative Research Environments. The model is based on experiences and research done in various Grid- and other collaborative projects. It was refined by interviews and discussion with developers, operators, and

users of CREs. In addition we drew on experiences and models of large software systems. Our interviews show that for CREs most efforts are currently spent on the prototyping phase, although experiences from development of large software systems show, that success of a product depends to a very large extent on early incorporation of constraints from operation and product transfer.

For the successful development of such a CRE it is crucial to consider all phases of the life cycle. This encompasses not only the developers and operators of such CRE, but also the funding partners or institutions, which have significant influence on the CRE. If the sponsors of CREs neglect to take the whole life cycle into account, the developers cannot be expected to resolve all connected issues by themselves.

We hope that our approach of combining general criteria and discipline specific elements into a life cycle model organizing the structural, functional, organizational, and financial aspects provides guidelines for all participants of CRE. These guidelines may help to build and operate successful CREs.

## References

- [CR10] Amy Carusi and Torsten Reimer. Virtual Research Environment Collaborative Landscape Study. <http://www.jisc.ac.uk/publications/reports/2010/vrelandscapestudy.aspx>, 2010. Accessed April 15, 2014.
- [dAIDI11] Arbeitsgruppe Virtuelle Forschungsumgebungen der Allianz-Initiative Digitale Information. Virtuellen Forschungsumgebungen - Ein Leitfaden. [http://www.allianzinitiative.de/fileadmin/user\\_upload/Leitfaden\\_VRE\\_de.pdf](http://www.allianzinitiative.de/fileadmin/user_upload/Leitfaden_VRE_de.pdf), 2011. Accessed April 15, 2014.
- [For11] DINI AG Virtuelle Forschungsumgebungen. Workshop: Virtuelle Forschungsumgebungen - erste Erfahrungen und Ergebnisse. <http://www.dini.de/veranstaltungen/workshops/forschungsumgebungen-2011/>, 2011. Accessed April 15, 2014.
- [For13a] Deutsche Forschungsgemeinschaft. Empfehlungen der Kommission "Selbstkontrolle in der Wissenschaft" - Vorschläge zur Sicherung guter wissenschaftlicher Praxis. [http://www.dfg.de/download/pdf/dfg\\_im\\_profil/reden\\_stellungnahmen/download/empfehlung\\_wiss\\_praxis\\_1310.pdf](http://www.dfg.de/download/pdf/dfg_im_profil/reden_stellungnahmen/download/empfehlung_wiss_praxis_1310.pdf), 2013. Accessed April 15, 2014.
- [For13b] Deutsche Forschungsgemeinschaft. RADIESCHEN - Rahmenbedingungen einer disziplinübergreifenden Forschungsdateninfrastruktur. Organisation und Struktur. [http://dx.doi.org/10.2312/RADIESCHEN\\_005](http://dx.doi.org/10.2312/RADIESCHEN_005), 2013. Accessed April 15, 2014.
- [LE13] Jens Ludwig and Harry Enke. Leitfaden zum Forschungsdaten-Management: Handreichungen aus dem WissGrid-Projekt. [http://www.wissgrid.de/publikationen/Leitfaden\\_Data-Management-WissGrid.pdf](http://www.wissgrid.de/publikationen/Leitfaden_Data-Management-WissGrid.pdf), 2013. Accessed April 15, 2014.
- [SUR12] SURFfoundation. Virtual Research Environments Starters Kit. <http://www.surf.nl/en/knowledge-and-innovation/knowledge-base/2011/vre-virtual-research-eenvironment-starters-kit.html>, 2012. Accessed April 15, 2014.

# Anforderungen eines nachhaltigen, disziplinübergreifenden Forschungsdaten-Repositoriums

Jan Potthoff<sup>1</sup>, Jos van Wezel<sup>1</sup>, Matthias Razum<sup>2</sup>, Marius Walk<sup>1</sup>

<sup>1</sup>Karlsruher Institut für Technologie (KIT)

<sup>2</sup>FIZ Karlsruhe - Leibniz-Institut für Informationsinfrastruktur

jan.potthoff@kit.edu, jos.vanwezel@kit.edu,

matthias.razum@fiz-karlsruhe.de, marius.walk@kit.edu

**Abstract:** Wissenschaftliche Erkenntnisse basieren zunehmend auf digitalen Daten, deren Publikation, Verfügbarkeit und Nachnutzung im Rahmen der guten wissenschaftlichen Praxis gewährleistet werden muss. Da auf die Daten in weiteren Arbeiten Bezug genommen werden kann und neue Erkenntnisse daraus entstehen, muss der Datenquelle ein besonderes Vertrauen entgegengebracht werden können. Das Projekt Research Data Repostitorium (RADAR) geht diese Herausforderung durch die Etablierung einer generischen Infrastruktur für die Archivierung und Publikation von Forschungsdaten an. Der Schwerpunkt wird auf Forschungsdaten aus Bereichen mit kleineren Datenmengen (small sciences) gelegt, da in diesen Forschungsdateninfrastrukturen meist noch fehlen. Die Nachhaltigkeit und Vertrauenswürdigkeit wird durch Bitstream Preservation, einer angestrebten Zertifizierung und einem sich selbst tragenden Betrieb mit einer Kombination aus Einmalzahlungen und institutionellen Angeboten gewährleistet.

## 1 Einleitung

Unter digitaler Langzeitarchivierung wird das Bewahren von digitalen Informationen über einen technologischen Wandel hinaus verstanden [Ne12]. Mit digitaler Information ist dabei jegliche Art von Information gemeint, die durch eine Sequenz von Bits (Bitstream) darstellgestellt werden kann. Die Information lässt sich im digitalen Bereich in drei Ebenen unterteilen. Die physische Ebene stellt die kodierte Information auf einem Speichermedium dar. Art und Weise hängt dabei vom verwendeten Speichermedium ab. Auf der logischen Ebene liegen die Informationen in einer Form vor, welche durch ein Informationssystem, anhand einer Formatspezifikation, die Semantik und Syntax beschreibt, interpretiert werden. Auf der konzeptuellen Ebene ist die Information schlussendlich für den Menschen interpretierbar. Eine Aufgabe der digitalen Langzeitarchivierung ist daher der Erhalt der drei Informationsebenen [Th02].

Neben der rein technischen Ausrichtung der digitalen Langzeitarchivierung sind des Weiteren organisatorische Prozesse zu berücksichtigen. Mit dem Open Archival Information System (OAIS) wird beispielsweise ein Referenzmodell beschrieben, dass nicht nur die technischen Komponenten eines solchen Systems berücksichtigt, sondern



auch Personen, die an den Prozessen, wie Einlagerung (Ingest), Lesen (Retrieval) und der Administration, beteiligt sind [CC12]. Es werden also sowohl Aufgabenbereiche und Verantwortlichkeiten definiert als auch Modelle der Funktionen und Informationen in einem digitalen Archiv. Durch die allgemeine Beschreibung der Datentypen und Funktionsabläufe wird das Referenzmodell als vom Einsatzbereich unabhängig angesehen und bildet so eine wesentliche Grundlage für Archivsysteme.

Sowohl die technischen als auch die organisatorischen Prozesse einer Infrastruktur zur digitalen Langzeitarchivierung müssen so entworfen sein, dass dem System Vertrauen entgegengebracht werden kann. Insbesondere ist dabei die Nachhaltigkeit zu adressieren.

## 2 Nachhaltigkeit und Vertrauenswürdigkeit

Um das Vertrauen in digitale Archive zu stärken, wurden im Rahmen des Projekts ArchiSafe, das im Rahmen der eGovernment-Initiative BundOnline 2005 an der Physikalisch-Technischen Bundesanstalt (PTB) durchgeführt wurde, Konzepte für ein vertrauenswürdiges elektronisches Archiv entworfen. Adressiert werden fünf Kernanforderungen: Vollständigkeit, Integrität und Authentizität, Lesbarkeit, Verfügbarkeit und Verkehrsfähigkeit [HR08]. Auch für den Bereich der Forschung wurden im Rahmen des von der Deutschen Forschungsgemeinschaft (DFG) geförderten Projekts "Beweissicheres elektronisches Laborbuch" (BeLab) Konzepte zur beweissicherhaltenden Archivierung entwickelt [Jo13]. Dabei wurden neben den technischen Möglichkeiten insbesondere juristische Fragestellungen berücksichtigt. Neben diesen Projekten wurden in den letzten Jahren Kriterienkataloge entworfen, anhand derer ein umgesetztes Archiv bewertet und zertifiziert werden kann, um das Vertrauen in die Infrastruktur zu fördern. Basierend auf [CR07] lassen sich die Anforderungen in drei Kategorien unterteilen: Organisatorische Strukturen, Datenmanagement und technische Aspekte.

### 2.1 Organisatorische Strukturen

Das *Anwendungsgebiet* des Systems und die Nutzergruppen sind zu definieren. Es müssen Leitlinien, die sowohl inhaltliche Kriterien als auch Betriebskriterien darlegen, zur allgemeinen Verfügung gestellt [SHH08] und Rechte und Pflichten des Anbieters (auch bezüglich Archivierungszeiträume) sowie der Urheber der Daten festgeschrieben werden [Ne10]. Das Ziel der *langfristigen Erhaltung* von Daten mit Möglichkeit der Verwaltung und des Zugriffs muss klar erkennbar sein [Mc12]. Dazu gehört auch die Definition von Strategien für das vorzeitige Ende des Betriebs oder das substantielle Ändern des Anwendungsbereiches [Do09]. Zur Erhaltung des Repositoriums müssen kurz- und langfristige Geschäftspläne (*Geschäftsmodell*) verfasst sein, die in einem zyklischen Prozess erarbeitet werden. Die Pläne müssen dabei mindestens jährlich überprüft und gegebenenfalls angepasst werden. Das Modell muss Einnahmen, Investitionen, Ausgaben und Risiken enthalten und dabei von Dritten transparent nach gesetzlichen Vorgaben überprüft werden können [RL02]. Die formale Prüfung und die Beurteilung des Repositoriums müssen regelmäßig erfolgen, so dass technologische

Entwicklungen und wachsende Anforderungen rechtzeitig erkannt werden. Daher ist eine regelmäßige *Zertifizierung* sinnvoll [CR07].

## 2.2 Datenmanagement und technische Aspekte

Das Gesamtsystem sollte *OAIS-konform* aufgebaut sein [SHH08]. OAIS bietet dabei ein allgemeines Framework aus Terminologien und Konzepten, um Architekturen und Operationen digitaler Archive zu beschreiben und zu vergleichen [Mc12]. Das Repositorium muss festlegen, welche *Eigenschaften eines digitalen Objekts* erhalten bleiben und in welcher Form es gespeichert wird. Beispielsweise kann zu einem Dokument nur der Text oder auch die Darstellung gesichert werden. Für die gängigsten Typen angenommener digitaler Objekte sollten Kriterien festgelegt sein, die genau beschreiben, um welche Objekte es sich handelt. Dem Objekt sollte dann ein *Persistent Identifier* zugewiesen werden [Do09]. Dabei gelten inhaltlich veränderte Dokumente als neue Dokumente. Die *Mindestzeit der Verfügbarkeit* eines Objekts für Dokumenten- und Publikationsservices darf beispielsweise 5 Jahre nicht unterschreiten und sollte in den Leitlinien festgehalten werden. Dabei sollte stets die Verbindung von Metadaten mit deren Objekte gewährleistet sein [Ne10].

Um alle Anfragen der Nutzer nach hinterlegten Daten erfüllen zu können, muss für jedes Objekt ein *Minimum an Metadaten*, die zur Identifizierung des Objekts dienen, hinterlegt sein. Hierzu muss das Repositorium schon von Beginn an eine Mindestmenge an Metadaten fordern [SHH08]. Das Repositorium sollte zudem aufzeigen, wie Metadaten gewonnen werden [CR07]. Es sollte die Möglichkeit bestehen über eine *Schnittstelle* standardisierte Metadaten ausgeben zu können [Ne10]. Die *Erreichbarkeit* des gesamten Angebots sollte über eine Webseite gewährleistet sein. Generell sollten aber verschiedene Möglichkeiten des Zugriffs vorhanden und frei zugänglich sein [CR07].

Das *Betriebskonzept* muss eine angemessene Verfügbarkeit des Systems garantieren und ausreichend dokumentiert sein. Die Dokumentation sollte Angaben über Komponenten, Zugangsregeln (personelle Verantwortlichkeit und Vertretung) und die pflichtmäßig durchzuführenden regelmäßigen Wartungen enthalten. Dabei sollten Technologien zur Sicherung und Wiederherstellung der Server-Software, der Metadaten und der Objekte eingesetzt werden [Ne10]. Das Repositorium muss dazu in der Lage sein, die Anzahl vorhandener Kopien digitaler Objekte und deren Ablageort zu definieren [B110].

## 3 Archivierungsinfrastrukturen

National aber auch international gibt es eine Reihe von Beispielen für Forschungsdaten-Repositorien, wie z. B. das World Data System der International Council of Science (ICSU)<sup>1</sup>, das Studien und empirischen Primärdaten aus den Sozialwissenschaften archiviert und zur Verfügung stellt. Im Bereich der Geowissenschaften ist ein weiteres Beispiel das Repositorium PANGAEA<sup>2</sup>. Aber auch Fachdatenbanken, wie z. B. die

---

<sup>1</sup> <http://www.icsu-wds.org/>

<sup>2</sup> <http://pangaea.de>

Protein Data Bank (PDB)<sup>3</sup>, die Informationen über experimentell gewonnene Daten von Proteinstrukturen und Nukleinsäuren zur Verfügung stellt, werden genutzt.

Die Verlässlichkeit der Technologie ist eine der wichtigsten Faktoren bei der Archivierung von Daten, die über mehrere Jahrzehnte gespeichert, lesbar und zugänglich sein müssen. In Rechenzentren werden Daten vor allem auf magnetischen Datenträgern in Form von Festplatten oder auf bis zu 700 Metern langen und 0,5 Zoll breiten Magnetbändern gespeichert. Beide Speichertechnologien haben eine vergleichbare Nutzungsdauer von vier bis sechs Jahren. Diese ist jedoch abhängig von der Anzahl der Lese- und Schreib-Vorgängen. Obwohl diese in einem Datenarchiv erwartungsgemäß gering sind, müssen die Speichermedien regelmäßig zur Überprüfung der Datenintegrität (Bitstream Preservation) gelesen werden. Am Ende der Lebensdauer einer Festplatte oder eines Bandes müssen die darauf befindlichen Daten ein letztes Mal zur Migration auf ein neues Gerät oder Medium gelesen werden.

Für große Archive und Langzeitarchive werden aus Kostengründen Magnetbänder zur Datenspeicherung eingesetzt. Bandspeichermedien sind nicht verbunden mit den eigentlichen Lese-/Schreibgeräten. Dadurch können sie mit einem niedrigeren Energieverbrauch betrieben werden und bieten so gegenüber Festplatten einen Kostenvorteil. Mit modernen Bandlaufwerken wird eine Abwärtskompatibilität bis zu drei Generationen gewährleistet, so dass in der Praxis eine Lebensdauer von acht bis zwölf Jahren erreicht werden kann. Ein weiterer Vorteil ist die in den letzten 10 Jahren angestiegene Speicherdichte und die damit um das 30- bis 50-fach angestiegene Speicherkapazität.

Systeme, auf denen Repositorien für eine verlässliche Archivierung aufbauen, bestehen aus zwei Kernkomponenten. Eine Datenbank, in der der Speicherort der archivierten Daten nach dem Ingest vermerkt wird, und ein System, das den Zugriff auf unterschiedliche Speichermedien abstrahiert. Zusätzlich sorgt es für die Migration der gespeicherten Daten zwischen den Speicherebenen oder Hierarchien auf der Grundlage von dem Benutzer definierter Regeln. Die Datenbank erlaubt des Weiteren die Ablage von Metadaten in unterschiedlichen Formaten. Kommerzielle Systeme, wie HPSS<sup>4</sup>, SAM-QFS<sup>5</sup> oder DMF<sup>6</sup>, sind in der Lage Millionen Objekte und mehrere Petabyte auf unterschiedlichen Speichersystemen, von Solid-State-Drive (SSD) bis Band, zu verwalten. Neben der Softwareentwicklung zum Datenmanagement in der Teilchenphysik entstanden die mit den kommerziellen Produkten vergleichbaren Open Source Pakete, wie zum Beispiel dCache und xrootd, die ebenfalls mit Petabyte-großen Speichersysteme umgehen und die eine automatische Migration zwischen Speicherebenen/-hierarchien unterstützen. iRODS, ein weiteres Open Source Paket, bietet ein regelbasiertes Datenmanagement mit dem eine automatische Migration in Abhängigkeit von den gespeicherten Metadaten, wie Größe oder Eigentümer, gesteuert werden kann.

---

<sup>3</sup> <http://www.rcsb.org/>

<sup>4</sup> <http://www.hpss-collaboration.org/>

<sup>5</sup> <http://www.oracle.com>

<sup>6</sup> <http://www.sgi.com>

## 4 Disziplinübergreifendes Forschungsdaten-Repository

Die Nachvollziehbarkeit und Reproduzierbarkeit von wissenschaftlichen Ergebnissen, etwa in [LMS12] unter dem Begriff *reproducible research* beschrieben, sowie neuartige Ansätze des Erkenntnisgewinns, die von Jim Gray als Fourth Paradigm bezeichnet wurden [HTT09], führen zu einer steigenden Bedeutung der Forschungsdaten im Rahmen von Veröffentlichungen. Verbunden damit ist der Bedarf, Forschungsdaten, auf denen eigene Ergebnisse beruhen, zu referenzieren. Das Referenzieren – in Form von Zitaten ein Kernelement wissenschaftlicher Kommunikation und Reputationsbildung – bedarf aber ebenso wie die Reproduzierbarkeit von wissenschaftlichen Erkenntnissen einen zuverlässigen und dauerhaften Zugriff auf die zugrunde liegenden Forschungsdaten und damit einer entsprechenden Infrastruktur.

### 4.1 Anwendungsbereich

Das von der DFG geförderte Projekt Research Data Repository (RADAR) stellt eine solche Infrastruktur für die Archivierung und Publikation von Forschungsdaten bereit. Partner des Projekts sind das FIZ Karlsruhe – Leibniz-Institut für Informationsinfrastruktur, die Technische Informationsbibliothek (TIB) in Hannover, das Steinbuch Centre for Computing (SCC) des KIT und zwei wissenschaftliche Partner (das Department für Chemie der LMU München und dem Leibniz-Institut für Pflanzenbiochemie in Halle).

Das Projekt richtet sich an Wissenschaftler (z. B. in Projekten) und Institutionen (z. B. Bibliotheken) mit zwei Angeboten: Ein disziplinübergreifendes Einstiegsangebot zur formatunabhängigen Datenarchivierung mit minimalen Metadatensatz und einem darauf aufbauenden Angebot mit integrierter Datenpublikation. Das Angebot der Archivierung (Einstiegsangebot) bietet eine formatunabhängige Archivierung mit Bitstream Preservation und einem minimalen Metadatensatz gemäß dem DataCite Metadata Kernel aus. Darüber hinaus erlaubt es eine Verknüpfung von Daten mit Metadaten.. Im Einstiegsangebot behält der Datengeber das alleinige Zugriffsrecht auf seine in RADAR archivierten Forschungsdaten, welches auf einen ausgewählten Benutzerkreis erweiterbar ist. Die so archivierten Forschungsdaten erhalten einen persistenten Identifier in Form eines Handles. Dieses Angebot richtet sich an Kunden, die primär an einer Datenarchivierung unter Einhaltung der von der DFG empfohlenen Haltefristen von 10 Jahren interessiert sind. Daneben können vom Benutzer weitere Haltefristen gewählt werden, die eine Mindesthaltefrist von 5 Jahren, siehe Abschnitt 2.2, nicht unterschreiten. Darüber hinaus eignet es sich aber auch für andere Daten wie z. B. Negativdaten<sup>7</sup>, die man nicht als Teil einer Publikation verwenden möchte. Das zweite Angebot unterstützt die dauerhafte Datenarchivierung mit integrierter Datenpublikation. Über die Leistungen des Einstiegsangebots hinaus ist hier die Vergabe von format- und disziplinspezifischen Metadaten sowie die Vergabe von dauerhaften DOI-Namen im Rahmen eines Publikationsprozesses vorgesehen. Durch diese Maßnahmen ist das Referenzieren von Daten, die zu einer Veröffentlichung im klassischen Sinne gehören,

---

<sup>7</sup> Daten, die zu keinem Ergebnis geführt haben.

aber auch eine reine Datenpublikation mit eigener DOI-Registrierung möglich. Dies bedingt auch eine aufwändigere Beschreibung der so publizierten und archivierten Forschungsdaten mit disziplinspezifischen Metadaten.

Gegenüber kommerziellen Diensten wie z. B. Dropbox<sup>8</sup> oder figshare<sup>9</sup> setzt sich RADAR in drei wesentlichen Punkten ab: Hinter RADAR steckt ein Betreiberkonsortium aus wissenschaftsnahen Gedächtnisorganisationen und Infrastruktureinrichtungen, die dem deutschem Recht unterliegen und keine kommerziellen Interessen verfolgen. Gegenüber kommerziellen Anbietern wird z. B. durch die Überwachung von Haltefristen und die Vergabe von persistenten Identifikatoren (Handles und DOIs) ein Mehrwert geboten. Drittens unterscheidet sich RADAR durch die zugesicherte Bitstream Preservation.

Als Entwicklungsgrundlage der RADAR Infrastruktur dient eine zuvor durchgeführte Anforderungsanalyse, in der insbesondere der Forschungsprozess und der zugehörige Workflow durch zwei an dem Projekt beteiligte Partner betrachtet wird. Dabei soll herausgestellt werden, welche Daten als erhaltungswürdig eingestuft werden können, welche Software in dem entsprechenden Forschungsbereich zum Einsatz kommt, und welche Datenformate daraus resultieren. Das Ziel der Analyse ist es, einen prototypischen Workflow zu genieren, anhand dessen die Anforderungen an die Archivierung im Allgemeinen, aber auch an einen forschungsnahen Daten Ingest, herausgestellt werden können. Basierend auf den beschriebenen zwei Angeboten sollen zudem im Rahmen der Anforderungsanalyse allgemeine und fachspezifische Metadatenfelder identifiziert werden. Bei der Ausgestaltung dieser Metadatenprofile kann über eine Kooperation mit dem Projekt Large Scale Data Management & Analysis (LSDMA)<sup>10</sup> am KIT auf umfangreiche Vorarbeiten zurückgegriffen werden. Thematisch legt RADAR den Schwerpunkt auf die *small sciences*, in denen Forschungsdateninfrastrukturen meist noch fehlen. RADAR erlaubt eine temporäre oder im Falle einer Datenpublikation zeitlich unbegrenzte Datenarchivierung. Das Projekt will nicht in Konkurrenz zu anderen, wie die in Abschnitt 3 genannten Beispiele, treten. Wie zuvor dargestellt sind die in den bisherigen Datenzentren vorhandenen Lösungen meist stark auf eine Zielgruppe ausgerichtet oder stellen kleinere Fachanwendungen für eine begrenzte Nutzergruppe dar. Hier versucht RADAR über das Angebot eines generischen, d. h. disziplinunabhängigen Repositoriums, die Lücke zu schließen und gleichzeitig über seine offenen Schnittstellen eine Basis für die Etablierung weiterer Fachanwendungen zu bilden.

## 4.2 Technische Aspekte

Als Frontend wird dem Nutzer eine intuitiv zu bedienende Web Oberfläche geboten. Hierüber lassen sich sowohl Daten und Metadaten erfassen als auch Regelungen, wie Servicelevel und Zugriffsrechte definieren. Zur schnelleren Übertragung von großen Datenmengen können des Weiteren direkte Speicherschnittstellen, die entsprechend

---

<sup>8</sup> <https://www.dropbox.com/>

<sup>9</sup> <http://figshare.com/>

<sup>10</sup> <http://www.helmholtz-lsdma.de/>

geeignete Protokolle unterstützen, genutzt werden. Die Gesamtarchitektur ist in Abbildung 1 dargestellt. Sowohl die im Abschnitt 3 beschriebenen kommerziellen als auch die Open Source Systeme bieten Funktionen zur Zugriffskontrolle. Des Weiteren unterstützen sie die Überprüfung der Datenintegrität auf der Grundlage von Prüfsummen. Die Absicherung der Datenintegrität kann durch das Anlegen von mehreren Kopien erhöht werden. So besteht die Möglichkeit mehrere Kopien der Daten anzufertigen und bei Bedarf und nach Ausstattung der Anlage auf unterschiedliche geografische Standorte zu verteilen. Dabei – insbesondere durch die Nutzung von Bändern als Speichermedium – muss mit einer längeren Zugriffszeit gerechnet werden. Möglichkeiten der Migration und schnellere Zugriffszeiten können durch den Betreiber durch unterschiedliche Servicelevel definiert werden, so dass der Nutzer selbst über Sicherheit der Datenintegrität und Zugriffszeiten in Abhängigkeit von entstehenden Kosten entscheiden kann. Die Bitstream Preservation ist nur eine Herausforderung der digitalen Langzeitarchivierung, auf die im Projekt RADAR besonderen Fokus gelegt wird. Teilweise existieren für weitere Aufgaben der Langzeitarchivierung bereits Lösungen oder müssen noch erforscht werden, z. B. im Bereich der Interpretation von Daten. Für diese soll das System über geeignete Schnittstellen flexibel sein.

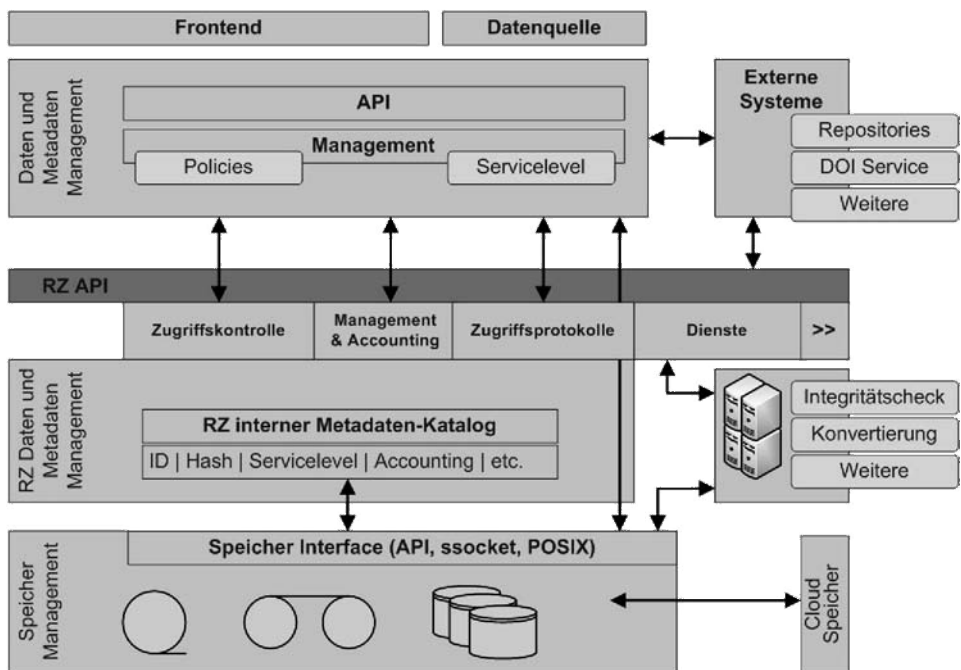


Abbildung 1: RADAR Architekturentwurf

Wie in Abbildung 1 dargestellt, ist der Kern der RADAR Architektur eine erweiterbare Schnittstelle für die Archiv Storage Services (Rechenzentrum Programmierschnittstelle – RZ API), die von einem Rechenzentrum angeboten werden. Die Schnittstelle verbirgt die Nutzung verschiedener Storage Services vor dem Frontend. Die RZ API stellt neben Schnittstellen für die Speicherung von Dateien Mechanismen für Zugangskontrolle,

Management, Accounting und zukünftig weitere inhaltsabhängige Services zur Verfügung. Da auf dieser Ebene auch die Bitstream Preservation sichergestellt werden muss, bietet sie die Möglichkeit der Neuberechnung von Prüfsummen. Dies ist aufgrund ansteigender Datenmengen nur über einen automatisierten Workflow möglich. Ein Automatismus könnte beispielsweise auch für Formatkonvertierungen umgesetzt werden. Vorgesehen ist, dass die RZ API auch Zugriffe auf das Metadaten System des Frontend ermöglicht und damit schnell Informationen über die gespeicherten Daten abfragen kann. Die Schicht oberhalb der RZ API übernimmt das Management von Kopien der Daten an unterschiedlichen Standorten und deren Erhalt. Außerdem übernimmt sie Koordinierungsaufgaben, die z. B. für die Vergabe von DOIs notwendig sind. Ein wichtiges Feature ist aber auch, dass die eingestellten Daten ohne Frontend abgerufen werden können. Dazu ist ein selbstbeschreibendes Dateiformat vorgesehen aus dem die Metadaten wieder hergestellt werden können.

Die Definition der RZ API soll es des Weiteren ermöglichen weitere Rechenzentren in das System einbinden zu können. Die RZ API Schicht trennt so Art der Speicherstrukturen des Rechenzentrums und Arbeitsweise des Daten und Metadaten Managements. Über die Anbindung weiterer Rechenzentren lässt sich der Erhalt der Datenintegrität verteilt über geografisch unterschiedliche Standorte realisieren. Des Weiteren können Benutzer die Standorte, an denen die Daten hinterlegt werden sollen, wählen. Neben der Anbindung von universitären Rechenzentren soll es möglich sein, auch kommerzielle Anbieter mit einzubinden. Für diese Möglichkeiten müssen entsprechende Services definiert und entsprechende Kosten im Rahmen des Geschäftsmodells kalkuliert werden.

Um das Vertrauen in das System zu stärken, soll eine Zertifizierung durchgeführt werden. Da eine große Anzahl von Kriterienkatalogen und Zertifizierungen existieren, muss in einem ersten Schritt eine Bewertung der Zertifikate und eine Eignungsprüfung für das angestrebte System durchgeführt werden. Die Zertifizierung wird sowohl das Datenzentrum als auch das darauf aufsetzende System (Management-Schicht) umfassen. Darüber hinaus soll auch das Verfahren der Datenpublikation zertifiziert werden.

### **4.3 Geschäftsmodell**

Ein wesentlicher Punkt zur Nachhaltigkeit des Projekts ist das selbsttragende Geschäftsmodell, das im Laufe des Projekts ausgearbeitet werden soll. Als Zielgruppe werden neben Forschungsprojekten (d. h. Wissenschaftlern) auch institutionelle Nutzer, wie z. B. Bibliotheken, die das Forschungsdatenmanagement in ihrer Institution selbst übernehmen, die Aufgabe der Bitstream Preservation aber an einen Dienstleister auslagern möchten, adressiert. Insbesondere Wissenschaftler sind jedoch nicht in der Lage nach Ende ihres Projekts dauerhaft für die Archivierung ihrer Daten zu bezahlen. Daher sieht das Geschäftsmodell Einmalzahlungen für die Archivierung und optionale Publikation der Forschungsdaten in Abhängigkeit von Datenumfang und Haltedauer vor. So können schon während der Antragstellung Kosten zur Archivierung und Datenpublikation abgeschätzt und mit beim Fördergeber beantragt werden.

Da RADAR (vorerst) nur die Bitstream Preservation sicherstellt und die funktionale Langzeitarchivierung (Content Preservation) mittels Formatüberwachungen, Migration und Emulation außen vor lässt, lassen sich die zu erwartenden Kosten für die Zukunft recht gut abschätzen. Die höchsten Kosten fallen beim Ingest der Daten (also einer initialen Aufbereitung der Daten mit Beratung, gegebenenfalls einmaliger Formatkonvertierung in ein geeignetes Archivformat, Anreicherung mit Metadaten usw.) einmalig an. Die Kosten für die Speicherung sind abhängig von der Datenmenge, nicht deren Format oder Komplexität. Die Speicherung von Daten wird – bei gleichbleibender Menge – seit Jahrzehnten günstiger. Das weltweit rasant ansteigende Datenvolumen [Ec10] lässt erwarten, dass sich dieser Trend auch in Zukunft fortsetzt, auch wenn dies gegebenenfalls auch einen Technologiewechsel erfordern wird. Ist der grundlegende Betrieb des Rechenzentrums, in dem die zu archivierenden Daten gespeichert sind, durch andere Aufgaben gewährleistet, werden die Kosten für die weitere Archivierung mit zunehmender Zeit marginalisiert [BCL08]. Durch den Entwurf der RZ API, siehe Abschnitt 4.2, soll gewährleistet werden, dass auch weitere Rechenzentren, die diese API zur Verfügung stellen, eine Archivierung der Daten durchführen können. Dadurch können Verantwortlichkeiten auf andere Anbieter ausgelagert werden.

## **5 Fazit und Ausblick**

Wie dargestellt, existiert eine Vielzahl von Richtlinien, Empfehlungen und Kriterienkatalogen, die den Aufbau eines vertrauenswürdigen Repositoriums darstellen und anhand derer sich eine Entwicklung bewerten lässt. Aufgabe wird es sein geeignete Schnittstellen zwischen Speicherschicht und Anwendungsschicht zu definieren. Die Herausforderung ist dabei insbesondere den Aufwand zur Implementierung der Schnittstelle so gering wie möglich zu halten, so dass Rechenzentren mit unterschiedlichsten Speicherstrukturen – aus Langzeitarchivierungssicht auch erwünscht – diese implementieren können. Bei der darüber liegenden Management-Schicht soll ebenfalls das Ziel einer möglichst flexiblen Schnittstelle verfolgt werden, um auch hier die Möglichkeit der Anbindung unterschiedlichster Frontends zu gewährleisten.

Datenspeicherung und vor allem die Datenarchivierung von Forschungsdaten ist eine auf Dauerhaftigkeit angelegte Dienstleistung. Der Großteil der dabei entstehenden Kosten soll auf die wissenschaftlichen Nutzer des Datenzentrums umgelegt werden, wobei bei der Entwicklung des Geschäftsmodells die aktuellen Gegebenheiten des Wissenschaftsbetriebs berücksichtigt werden müssen. Hier existiert eine starke Ausrichtung auf geförderte Projekte mit begrenzter Laufzeit. Nach dem Ende des Projekts besteht daher keine Möglichkeit weiterhin für die Datenarchivierung zu bezahlen. Eine direkte Förderung des Datenzentrums wird bei dem angedachten Geschäftsmodell weitgehend vermieden. Die Ausarbeitung eines Kostenmodells, das die benutzten Speichermodalitäten und zukünftige technische Entwicklungen berücksichtigt, muss erarbeitet werden. Darüber hinaus müssen vertragliche und rechtliche Regelungen für den Betrieb und das Dienstleistungsangebot des Datenzentrums erarbeitet werden. Wichtige Aspekte sind dabei die Weiterverwertung der Daten durch Dritte, Haftungsfragen und die Dauer der Datenspeicherung.



## Literaturverzeichnis

- [BCL08] Beagrie, N.; Chruszcz, J.; Lavoie, B.: Keeping Research Data Safe - A Cost Model and Guidance for UK Universities. Final Report. s.l.: JISC, 2008, S. 91-92, <http://www.jisc.ac.uk/media/documents/publications/keepingresearchdatasafe0408.pdf>, abgerufen am 10.04.2014.
- [B110] Blue Ribbon Task Force: Sustainable economics for a digital planet: Ensuring long-term access to digital information. In: Final Report of the Blue Ribbon Task Force on Sustainable Digital Preservation and Access, 2010, [http://brtf.sdsc.edu/biblio/BRTF\\_Final\\_Report.pdf](http://brtf.sdsc.edu/biblio/BRTF_Final_Report.pdf), abgerufen am 10.04.2014.
- [CC12] CCSDS: Reference Model for an Open Archival Information System (OAIS). Magenta book. Technischer Bericht, CCSDS, 2012.
- [CR07] Center for Research Libraries (CRL): Trustworthy Repositories Audit & Certification: Criteria and Checklist (TRAC), Chicago, 2007, [http://www.crl.edu/sites/default/files/attachments/pages/trac\\_0.pdf](http://www.crl.edu/sites/default/files/attachments/pages/trac_0.pdf), abgerufen am 10.04.2014.
- [Do09] Dobratz, S. et al.: Catalogue of Criteria for Trusted Digital Repositories, nestor materials, Deutsche Nationalbibliothek, Frankfurt (Main), Germany, <http://nbn-resolving.de/urn:nbn:de:0008-2010030806>, abgerufen am 10.04.2014.
- [Ec10] The Economist Newspaper Limited: The Data Deluge. In: The Economist, 2010. <http://www.economist.com/node/15579717>, abgerufen am 10.04.2014.
- [HR08] Hackel, S.; Roßnagel, A.: Langfristige Aufbewahrung elektronischer Dokumente: In: Klumpp (Hrsg.), Informationelles Vertrauen für die Informationsgesellschaft, Berlin 2008.
- [HTT09] Hey, T.; Tansley, S.; Tolle, K.: The Fourth Paradigm: Data-Intensive Scientific Discovery. Redmond, Washington : Microsoft Research, 2009.
- [Jo13] Johannes, P.C.; Potthoff, J.; Roßnagel, A.; Neumair, B.; Madiesh, M.; Hackel, S.: Beweissicheres elektronisches Laborbuch – Anforderungen, Konzepte und Umsetzung zur langfristigen, beweiswerterhaltenden Archivierung elektronischer Forschungsdaten und –dokumentation. In: Roßnagel, A. (Hrsg.): Der Elektronische Rechtsverkehr, Band 29, NOMOS, 2013.
- [Mc12] McGovern, N.Y.: Aligning National Approaches to Digital Preservation. Atlanta, USA, 2012, [http://educopia.org/sites/educopia.org/files/ANADP\\_Educopia\\_2012.pdf](http://educopia.org/sites/educopia.org/files/ANADP_Educopia_2012.pdf), abgerufen am 10.04.2014.
- [Ne10] Netzwerkinformation e.v., Deutsche Initiative für und Arbeitsgruppe Elektronisches Publizieren: DINI-Zertifikat Dokumente- und Publikationservice 2010, <http://nbn-resolving.de/urn:nbn:de:kobv:11-100182794>, abgerufen am 10.04.2014.
- [Ne12] Neuroth, H.; Oßwald, A.; Scheffel, R.; Strathmann, S.; Ludwig, J.: Langzeitarchivierung von Forschungsdaten: eine Bestandsaufnahme. Hülbusch, Boizenburg, 2012. <http://nestor.sub.uni-goettingen.de/bestandsaufnahme/index.php>, abgerufen am 10.04.2014.
- [RL02] Research Libraries Group (RLG): Trusted Digital Repositories: Attributes and Responsibilities, Report, 2002. <http://www.oclc.org/content/dam/research/activities/trustedrep/repositories.pdf>, abgerufen am 10.04.2014.
- [SHH08] Sesink, L.; van Horik, R.; Harmsen, H.: Data Seal of Approval, Data Archiving and Networked Services (DANS), Den Haag, 2008, <http://www.datasealofapproval.org/>, abgerufen am 10.04.2014.
- [LMS12] LeVeque, R.J.; Mitchell, I.A.; Stodden, V.: Reproducible research for scientific computing: Tools and strategies for changing the culture. In: Computing in Science and Engineering. July/August 2012, Bd. 14, 4, S. 13-17.
- [Th02] Thibodeau, K.: Overview of technological approaches to digital preservation and challenges in coming years. The state of digital preservation: an international perspective, Seiten 4–31, 2002.

# **Langzeitarchivierung von Forschungsdaten im Sonderforschungsbereich 840 „von partikulären Nanosystemen zur Mesotechnologie“ an der Universität Bayreuth**

Johannes Fricke, Dr. Andreas Weber und Dr. Andreas Grandel

Universität Bayreuth  
IT-Servicezentrum  
Gebäude NW II  
Universitätsstraße 30  
95447 Bayreuth

johannes.fricke@uni-bayreuth.de  
andreas.weber@uni-bayreuth.de  
andreas.grandel@uni-bayreuth.de

**Abstract:** Im Rahmen des Infrastruktur-Teilprojekts Z2 des Sonderforschungsbereichs 840 (SFB840) hat das IT-Servicezentrum (ITS) der Universität Bayreuth die Aufgabe, ein Forschungsdatenportal zu konzipieren und bereitzustellen. In dem Forschungsdatenportal sollen alle forschungsrelevanten Daten, sowie die für eine Nachnutzung notwendigen Metadaten gespeichert und archiviert werden. Die eindeutige, persistente Referenzierbarkeit der Daten soll über zertifizierte DOIs (Digital Object Identifier) sichergestellt werden. Besonderes Augenmerk soll dabei auf die langfristige Verfügbarkeit der Daten gelenkt werden. Dabei soll auf existierende, möglichst freie Softwarelösungen zurückgegriffen werden, die an den spezifischen Bedarf des SFB angepasst werden. Diese Lösung soll später die Grundlage für die Speicherung von Forschungsdaten auch aus anderen Bereichen bilden.

## **1. Ausgangslage**

Die nachhaltige Sicherung und die langfristige Zugänglichkeit von Forschungsdaten gewinnen für die Wissenschaft zunehmend an Bedeutung. Die wichtigsten Gründe dafür sind die Forderung nach Transparenz der Forschungsergebnisse und das Potential, das in der Nachnutzung bereits vorhandener Ergebnisse gesehen wird. Der offene Zugang zu

den wissenschaftlichen Ergebnissen wird in vielen Erklärungen der Wissenschaft gefordert [Be10]. Die Deutsche Forschungsgemeinschaft verlangt deshalb seit 2010 bei der Beantragung von Projekten eine Darstellung des Umgangs mit den wissenschaftlichen Ergebnissen [DF11]. Im Rahmen eines Teilprojektes des Sonderforschungsbereiches 840 (SFB840) [Z2] hat das IT-Servicezentrum (ITS) der Universität Bayreuth die Aufgabe, Möglichkeiten für die langfristige Speicherung und Wiederauffindbarkeit der forschungsrelevanten Daten des SFB840 zu konzipieren und bereit zu stellen. Ein Forschungsdatenportal, in dem die relevanten Metadaten abgelegt und abrufbar sind, soll den Forschungsprozess der im SFB beteiligten Wissenschaftler unterstützen. Auf die langfristige Verfügbarkeit der Forschungsprimärdaten, vorwiegend Bild- und Messdaten, wird besonderes Augenmerk gelegt. Es soll somit gewährleistet werden, dass die Daten über einen längeren Zeitraum zur Verfügung stehen und nicht in den Speicherorten der Labors durch Umzug oder Personalveränderungen verloren gehen.

Der Zugriff auf die Forschungsdaten soll über einen weltweit eindeutigen Identifier erfolgen. Über diesen Identifier können die Daten auch mit den Publikationen verknüpft werden und umgekehrt. Die Forschungsdaten werden somit zukünftig auch bei der Aufnahme der Publikationen in den Nachweiskatalogen verzeichnet. Ebenfalls kann zu den Forschungsdaten abgelegt werden, in welchen Publikationen die Daten eingeflossen sind. Die Forschungsdaten können somit auch externen Interessenten (z. B. Referees, anderen Arbeitsgruppen) zugänglich gemacht werden.

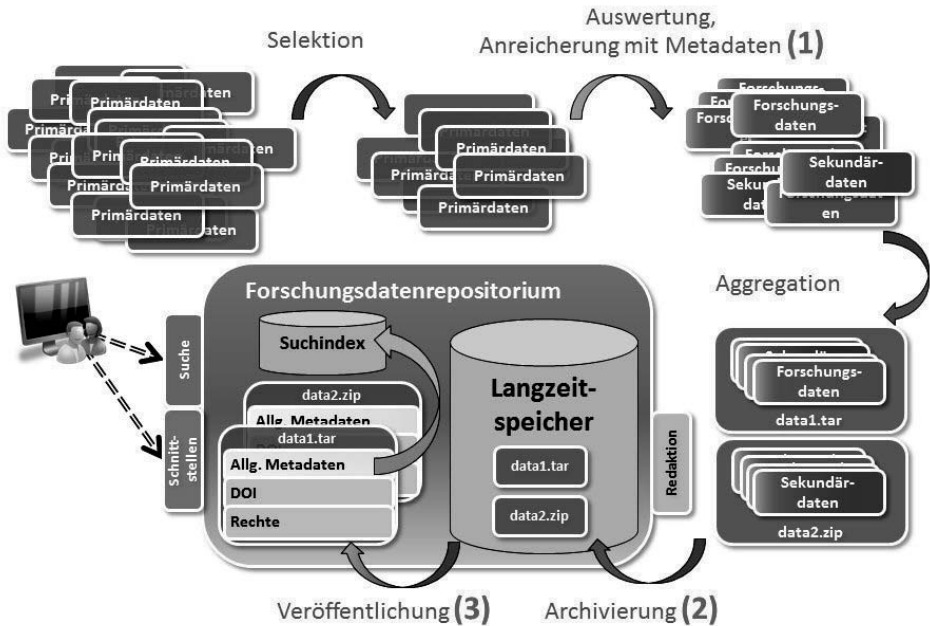
Die Beschreibung der Forschungsergebnisse mit geeigneten Metadaten ist ein wichtiger Bestandteil des Projektes, für den die Universitätsbibliothek als Partner gewonnen werden konnte. Das dort vorhandenen Wissen zu Metadaten und die langjährige Erfahrung in der Erschließung und Beschreibung von Informationsquellen soll genutzt werden, in Zusammenarbeit mit den Wissenschaftlern gute und praktikable Metadatenschemata zu entwickeln.

Der SFB840 überspannt viele Fachbereiche der Physik und der Chemie. Deshalb muss der Ansatz hinsichtlich der Art der gespeicherten Daten und der Beschreibungen der Daten flexibel sein. Die Erweiterbarkeit der Lösung auf die Archivierung von Forschungsdaten über den SFB hinaus sollte deshalb grundsätzlich möglich sein und für alle Wissenschaftler der Universität Bayreuth nutzbar sein.

## **2. Stand des Projekts**

### **2.1 Beschreibung der Aufgabe**

Die grundsätzlichen Komponenten für eine langfristige Speicherung von Forschungsdaten und deren Veröffentlichung ist in der, dem Antrag entnommenen, Abbildung 1 dargestellt. Die technische Umsetzung soll in diesem Papier diskutiert und konkretisiert werden.



**Abb. 1:** Notwendige Veränderung im Umgang mit Forschungsdaten: (1) Anreicherung mit allgemeinen und fachlichen Metadaten; (2) Transfer in ein Langzeitarchiv; (3) Veröffentlichung und Referenzierbarkeit

Erläuterungen zu der Abbildung 1:

- (1) Nach der Selektion der relevanten Daten werden die Primärdaten mit geeigneten Metadaten versehen, damit sie später lokalisiert und interpretiert werden können. Die allgemeinen Metadaten ermöglichen die Zuordnung der Daten zu einer Forschergruppe und einem Projekt und werden in dem Portal für das Auffinden und die Selektion der Datensätze verwendet. Die fachlichen Metadaten ermöglichen die wissenschaftliche Interpretation, Verifikation und Nachnutzung der Daten. Die Vergabe der Metadaten geht im Umfang über die Abbildung eines Laborbuches hinaus, da die Ergebnisse auch für Personen nachvollziehbar sein sollen, die den Versuchsaufbau, die verwendeten Programme und die verwendete Ausstattung nicht kennen. Die bei der Auswertung und Weiterverarbeitung der Primärdaten entstehenden „Sekundärdaten“ (Grafiken, Diagramme etc.) werden ebenfalls mit diesen Informationen versehen.
- (2) Die Daten verbleiben langfristig nicht mehr auf den Speichersystemen der Forschergruppen, sondern werden in einem hochverfügbaren und zuverlässigen Langzeitarchiv abgelegt. Damit wird sichergestellt, dass die Daten über mindestens 10 Jahre verfügbar bleiben, selbst wenn sich die Struktur oder die Organisation der Forschungsgruppen ändern (z. B. Weggang von Forschern etc.). Der Zugriff auf die Daten muss für die Forscher weiterhin jederzeit möglich sein. Im einfachsten Fall werden die Daten in ein Netzlaufwerk verschoben (ähnlich

zu Dropbox), aber auch die für die Übertragung von Daten im Internet üblichen Protokolle (SFTP, SCP, CIFS etc.) können verwendet werden. Im Projekt wird die Nutzung entfernter, zentraler Langzeitspeicher konzeptionell betrachtet.

- (3) Der Zugriff auf die Daten soll für externe Interessierte (z. B. Referees, andere Forschergruppen) möglich sein. Der Zugriff erfolgt über das Repository, das neben der Suchmöglichkeit in den allgemeinen Metadaten auch eine Prüfung der Zugriffsberechtigung bereitstellt. Letztere stellt sicher, dass die Daten erst nach der Veröffentlichung oder nach der Erteilung von Patenten für Externe zugänglich werden, vorab aber schon in Kooperationen (z. B. innerhalb des SFB) bereit stehen. Das Portal stellt somit auch eine Kollaborationsplattform für die Wissenschaftler dar. Die eindeutige, persistente Referenzierbarkeit der Daten wird über zertifizierte Identifikatoren erreicht. Hierzu soll der bereits seit 2005 etablierte DOI-Dienst der TIB Hannover genutzt werden, der im internationalen Kontext von DataCite errichtet wurde. Über die DOIs (Digital Object Identifier) können die Datensätze mit den Veröffentlichungen verknüpft werden und umgekehrt den Daten die Identifier der elektronischen Veröffentlichung zugeordnet werden. Über geeignete Schnittstellen werden die Metadaten für andere Systeme (z. B. übergreifenden Fachrepositorien) zugänglich gemacht.

## 2.2 Langzeitarchivierung

Seit Beginn des Projektes wurde insbesondere die Langfristspeicherung von Daten genauer betrachtet und eine Reihe von Veranstaltungen zu diesem Thema besucht [Bi13]. Die Evaluation ergab ein paar grundlegende Ergebnisse, die nachfolgend kurz erläutert werden.

Der OAIS-Standard (Open Archival Information System) [Spac] sollte bei der Langzeitarchivierung eingehalten werden. Das System zur Speicherung der Daten muss OAIS-konform sein, so dass verschiedene Lösungen darauf aufsetzen können. Damit wird auch eine Austauschbarkeit zwischen verschiedenen Systemen erreicht.

Wegen der Langfristigkeit der Aufgabe sollte bei der Langzeitarchivierung die Verwendung eigener Lösungen, deren Fortbestand an einzelne Personen geknüpft ist, vermieden werden. Es sollten am Markt etablierte Lösungen eingesetzt werden.

Die Beschreibung der Datensammlungen mit Metadaten sollte nach einem Standard, z.B. METS (Metadata Encoding & Transmission Standard) erfolgen [MET]. Die Metadaten müssen projektabhängig definiert werden können.

Die Sicherstellung der Integrität der Daten auf Bit-Ebene und automatische Formatkonversionen sollten durch das Speichersystem zur Verfügung gestellt werden. Die Verwaltung dieser Aufgaben sollte in dem System verankert sein.

Wegen der großen Datenmengen eignen sich einige der im Bereich der Verwaltung von digitalen Publikationen etablierte Open-Source Lösungen nicht, z.B. LOCKSS (Lots Of Copies Keep Stuff Safe) [LOC].

Die Daten sollen einfach zitiert werden können. Dazu sind geeignete Identifier zu erstellen.

Der Aufwand den die Forscher zur Aufbereitung der Daten für die Langzeitarchivierung betreiben müssen, insbesondere die Vergabe von Metadaten und der Transfer in das Archivsystem, muss gering sein. Deshalb müssen komfortable Werkzeuge für diesen Vorgang bereitgestellt werden.

Die Bereitschaft der Forscher, die Forschungsdaten als Open-Source bereit zu stellen ist sehr gering. Die Einsicht in die Daten soll aus Sicht der Forscher eher selektiv freigegeben werden. Deshalb ist ein ausgereiftes Rechtekonzept notwendig.

Die meisten Forscher kennen die Verpflichtungen zur Aufbewahrung der Daten nicht oder nicht ausreichend. Deshalb wird empfohlen, dass neben der technischen Umsetzung geeignete Richtlinien für den Umgang mit Forschungsdaten erstellt werden. Die Archivierung ist aus Sicht der Forscher nur für einen begrenzten Zeitraum notwendig, es besteht grundsätzlich nicht der Anspruch der „unendlich langen“ Bereithaltung.

Auf der Grundlage dieser Bedingungen wurde nach einem geeigneten System gesucht. Ein etabliertes Produkt im Bereich Langzeitspeicherung ist Rosetta [ExLi] von ExLibris, das in der Bayerischen Staatsbibliothek (betrieben vom Leibniz-Rechenzentrum) bereits im Bereich Digitalisierung im Einsatz ist. Rosetta bildet auch die Grundlage für die Langzeitspeicherung von Forschungsdaten an der ETH Zürich. Dort wird der Einsatz von Rosetta über das Problemfeld Forschungsdaten hinaus angestrebt, was in der Abbildung 2 aufgezeigt ist. Es wird versucht, die Bewahrung aller anfallenden digitalen Daten in einem Konzept zu vereinen und auf einer gemeinsamen Plattform umzusetzen. Die grundlegende Problematik bei der Langzeitarchivierung ist in den verschiedenen Bereichen gleich, die Vorgänge unterscheiden sich nur durch die Auswahl der Objekte und die Beschreibung mit Metadaten.

## VISION: ROSETTA ALS GEMEINSAME BASIS

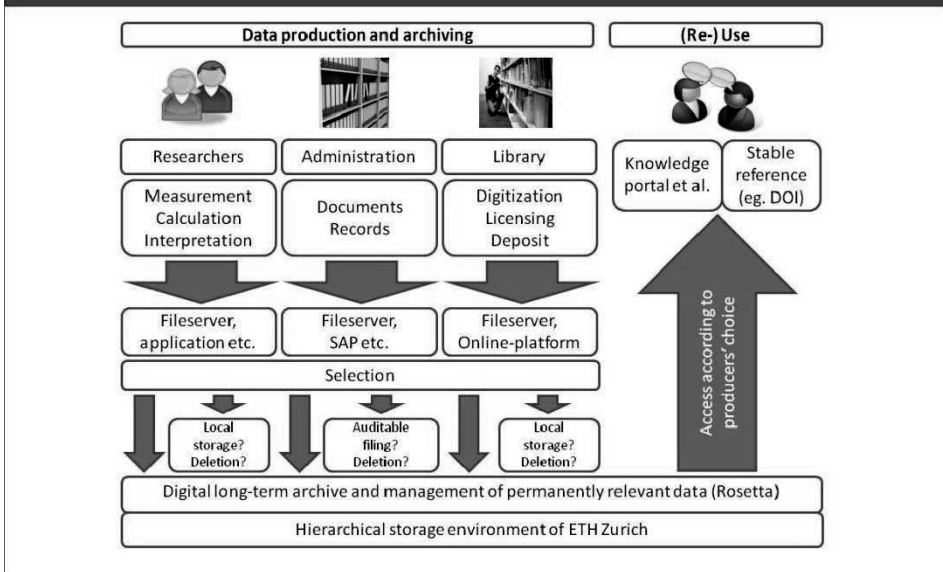
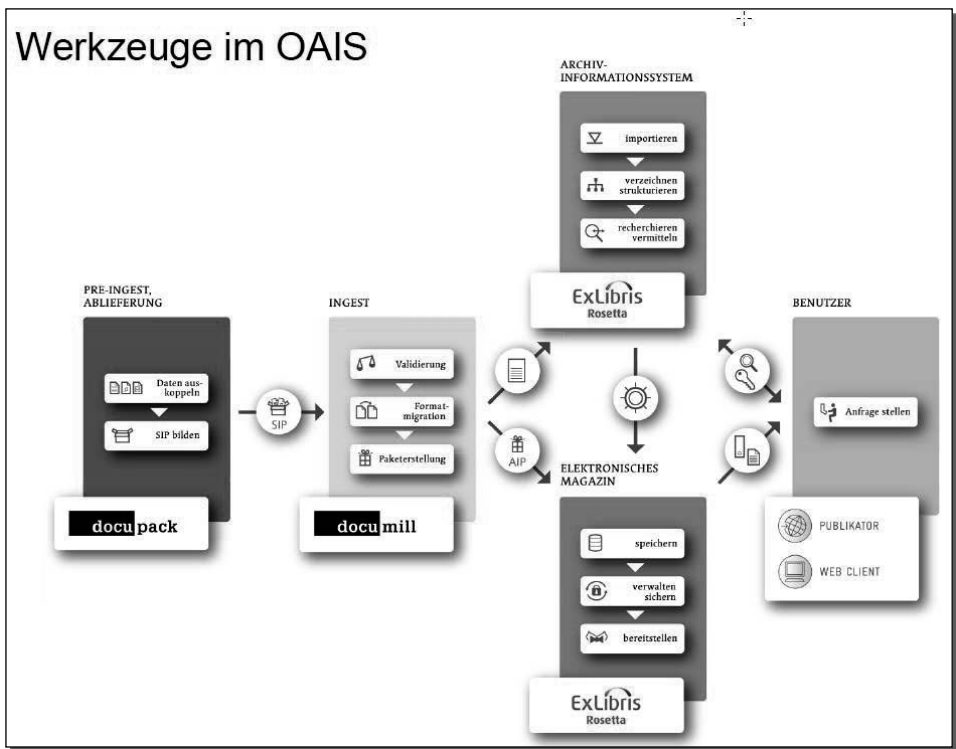


Abb. 2: **Ziel der ETH Zürich** (aus Töwe, Matthias. Forschungsdaten an der ETH Zürich, 2010 <http://dx.doi.org/10.3929/ethz-a-007590685>)

### 2.3 Transfer der Daten, Pflege der Metadaten

Zur Ablieferung der Daten durch die Forscher wird, an der ETH Zürich unter anderem auch die Software „Packer“ (ehemals docupack, vgl. Abb. 3) der Firma Docuteam [Docu] verwendet. Packer liefert standardisierte SIPs (Submission Information Package), woraus am Archivierungssystem durch das Programm Feeder (ehemals documill, vgl. Abb. 3) für Rosetta konforme AIPs (Archival Information Package) erzeugt werden. Das Zusammenspiel dieser Komponenten ist in der Abbildung 3 aufgezeigt.

Mit Packer werden die Aufgaben der Aggregation und die Anreicherung mit Metadaten, sowie der Transfer zum Langzeitarchiv abgedeckt. An dieser Stelle muss nun noch als externer Prozess die Erstellung einer DOI eingefügt werden. Die Einstellung der Daten und der Metadaten in das Archiv erfolgt über Feeder. Die Archivierung und die Bereitstellung der Daten erfolgt über Rosetta.



**Abb 3:** Eingesetzte Komponenten an der ETH Zürich (aus Strukturierung und Aufbereitung von Forschungsdaten fürs Archiv: Konzepte und Werkzeuge - Andreas Nef, Docuteam GmbH)

Diese Vorgehensweise stellt eine einfache Möglichkeit dar, Daten in das System zur Langfristspeicherung einzubringen. Eine Suche in den Metadaten oder eine nachträgliche Pflege der Metadaten wird damit aber noch nicht möglich.

Die Evaluation vorhandener Plattformen zur Erfassung und Verwaltung von Metadaten zu Forschungsdaten hat gezeigt, dass es bereits einige solcher Systeme gibt. Diese sind aber meistens für spezielle Anwendungsfälle ausgelegt. Die verwendeten Metadaten sind für alle Datensätze gleich und die Bearbeitung speziell dafür optimiert [DWB]. Eine Veränderung oder ganz flexible Gestaltung der Metadaten ist in diesen Systemen in der Regel nicht vorgesehen.

Das Problem der flexiblen Beschreibung von Objekten ist im Bereich der Bereitstellung von digitalen Sammlungen bereits bekannt. Auch hier ist die Beschreibung abhängig von den Objekten, die angeboten werden. Neben einer Reihe kommerzieller Programme gibt es die an der TU München entwickelte Software MediaTUM [Med]. Es handelt sich dabei um eine Open-Source-Software, die im Rahmen des DFG-Projekts IntegraTUM von der Universitätsbibliothek der TU München in enger Kooperation mit den Fakultäten entwickelt wurde. Die Metadaten sind hier frei definierbare und flexible Workflows für verschiedene Aufgaben erstellbar. Damit steht ein Portal zur Verfügung, in welchem



die Daten übersichtlich dargestellt, mit Vorschaubildern versehen werden und die Metadaten durchsucht und angezeigt werden können. Weiterhin ist eine Schnittstelle zur Anbindung eines Archivsystems bereits implementiert. Der hohe Grad an Modularisierung von mediaTUM erlaubt den einfachen Austausch einzelner Komponenten und damit scheint mit MediaTUM auch für die Beschreibung von Forschungsdaten geeignet. Werden auch die digitalen Sammlungen mit MediaTUM verwaltet, ergäbe sich damit zudem ein Synergieeffekt. Weiterhin ist zu erwähnen, dass die Universitätsbibliothek der Hochschule der Bundeswehr in Neubiberg, die Universitätsbibliothek der Universität Augsburg und die Katholische Universität Eichstätt bereits mediaTUM für die Verwaltung von digitalen Sammlungen einsetzen. Somit sollte die Weiterentwicklung langfristig gesichert sein und es sollten weitere Kooperationspartner zur Verfügung stehen.

### **3. Aussichten**

Die Evaluation der Komponenten, die als Basis für die Verwaltung von Forschungsdaten in Hinblick auf die Langfristspeicherung eingesetzt werden sollen, ist weitgehend abgeschlossen.

Die Bayerische Staatsregierung hat die Software Rosetta für die Nutzung durch die Bayerischen Hochschulen lizenziert [DiLa]. Für die langfristige Speicherung der Forschungsdaten des SFB 840 wird somit der Einsatz von Rosetta angestrebt. Zunächst wird hier auf die Installation in der BSB zurückgegriffen, ein späterer lokaler Einsatz von Rosetta wäre im Falle von Engpässen in den Ressourcen, z.B. Datenübertragung ebenfalls möglich. Der Umgang mit Rosetta wird derzeit genauer betrachtet.

Die Komponenten der Firma Docuteam sollen für den Transfer nach Rosetta eingehender evaluiert werden. Gespräche über eine Zusammenarbeit hierzu haben mit der BSB bereits stattgefunden. Mit diesen Komponenten wäre ein einfacher Transfer der Daten in den Langfristspeicher realisiert und damit eine einfache Basislösung gefunden.

Für den Aufbau eines Forschungsdatenportals wird der Einsatz von MediaTUM eingehend in einer lokalen Testinstallation der Software evaluiert. Gespräche mit den Entwicklern von MediaTUM haben bereits die Grundlage für eine mögliche Kooperation geschaffen. Eine wichtige Aufgabe wird es sein, MediaTUM und Rosetta in geeigneter Weise miteinander zu verbinden. Die Einbindung der DOI-Generierung ist ebenfalls zeitnah zu betrachten.

Die Erarbeitung von Metadatenschemata wird anhand einzelner Experimente exemplarisch durchgeführt. Hierzu sind intensive Gespräche mit den Wissenschaftlern notwendig. Ebenfalls wird versucht, einfache Werkzeuge für einen automatisierbaren Transfer der Metadaten zu entwickeln. Die Speicherung der Daten von Spektrometern soll dabei die Pilotanwendung darstellen.

## Literaturverzeichnis

- [Be10] Berliner Erklärung über den offenen Zugang zu wissenschaftlichem Wissen, 2003; OECD Declaration on Access to Research Data from Public Funding, 2004, Allianz der deutschen Wissenschaftsorganisationen, Grundsätze zum Umgang mit Forschungsdaten, 2010
- [Bi13] Bibliothekartag 2013, Session Forschungsdaten-Repositorien - Infrastrukturen zur dauerhaften Zugänglichkeit von Forschungsdaten, 11.03.2013, <http://www.bid-kongress-leipzig.de/t3/index.php?id=26>;  
Tagung: Digitale Langzeitarchivierung an Hochschulen, <http://www.ub.hu-berlin.de/de/ueber-uns/veranstaltungen/tagung-digitale-langzeitarchivierung-an-hochschulen>;  
Treffen der INF-Projekte, 11.4.2013, Göttingen
- [DF11] DFG-Antrag: Leitfaden für die Antragstellung, 2011. [http://www.dfg.de/formulare/54\\_01/54\\_01\\_de.pdf](http://www.dfg.de/formulare/54_01/54_01_de.pdf).
- [DiLa] Pressemitteilung „Digitales Langzeitgedächtnis für die bayerischen Hochschulen: <http://www.bayern.de/Pressemitteilungen-.1255.10476313/index.htm>
- [Docu] Hauptseite der Firma Docuteam GmbH aus Baden-Dättwil, Schweiz. Motto: „Informationen strukturieren und erhalten“: <http://www.docuteam.ch/>
- [DWB] Hauptseite des Deutschen Klimarechenzentrums zum Thema Langzeitarchivierung: [http://www.dkrz.de/daten/langzeitarchivierung?set\\_language=de](http://www.dkrz.de/daten/langzeitarchivierung?set_language=de);  
Hauptseite der Diversity Workbench, einer virtuellen Forschungsumgebung für Biodiversität: [http://diversityworkbench.net/Portal/Main\\_Page](http://diversityworkbench.net/Portal/Main_Page)
- [ExLi] Internetseite der ExLibris Gruppe über Rosetta / Langzeitarchivierung unter der Überschrift „Eine neue Methode zum Erhalt des kulturellen Erbes und kumulierten Wissens“: <http://www.exlibrisgroup.com/de/default.asp?catid=%7B12040B74-9794-451D-9369-28C0F2980F1B%7D>
- [LOC] Hauptseite von „Lots of Copies Keep Stuff Safe“ des LOCKSS Programms, welches in Zusammenarbeit mit den Stanford Universitätsbibliotheken betrieben wird: <http://www.lockss.org/>
- [Med] Hauptseite von mediaTUM: <https://mediatum.ub.tum.de/>
- [MET] Hauptseite des METS-Standards der Library of Congress (LOC): <http://www.loc.gov/standards/mets>
- [Spac] Space data and information transfer systems -- Open archival information system (OAIS) -- Reference el [http://www.iso.org/iso/home/store/catalogue\\_tc/catalogue\\_detail.htm?csnumber=57284](http://www.iso.org/iso/home/store/catalogue_tc/catalogue_detail.htm?csnumber=57284)
- [Z2] Projekt Z2 im Fortsetzungsantrag zum Sonderforschungsbereich 840 „von partikulären Nanosystemen zur Mesotechnologie“ <http://www.sfb840.uni-bayreuth.de/de/z-projects/index.html>



# **Energieeffizienz, Green-IT**



# Energieeffizienz im WLAN

Maximilian Boehner, Sebastian Porombka, Gudrun Oevel

Universität Paderborn  
Warburger Str. 100  
33098 Paderborn  
mboehner@mail.upb.de  
sebastian.porombka@upb.de  
gudrun.oevel@upb.de

**Abstract:** Obwohl der Energieverbrauch in WLAN-Infrastrukturen weltweit steigt, sind nur wenige Ansätze zur Optimierung ihrer Energieeffizienz bekannt. Ausgehend von der WLAN-Installation an der Universität Paderborn, wird zunächst eine genaue Analyse des IST-Zustands vorgenommen, um Potentiale zu entdecken und Metriken zu definieren. Anschließend werden verschiedene Strategien zur Senkung des Energieverbrauchs diskutiert. Zwei neu entwickelte Algorithmen zum Resource-on-Demand-Betrieb von Access Points werden vorgestellt und bewertet. Insgesamt zeigt sich, dass erhebliche Energieeinsparungen möglich sind, ohne signifikant auf Dienstgüte verzichten zu müssen.

## 1 Einleitung

GreenIT als Schlagwort kennzeichnet für große Organisationen die Herausforderung, einerseits stark wachsende IT-Infrastrukturen bereitzustellen, andererseits aber auch die Energiekosten im Auge zu behalten. Dies gilt insbesondere auch für den kabellosen Zugang zum Netzwerk (Funknetz, WLAN). Größere WLAN-Infrastrukturen werden wie viele IT-Komponenten deutlich überprovisioniert, um ausreichend Reserven für Leistungsspitzen vorzuhalten. Diese werden aber nur selten benötigt ([Ja07]), sodass die meiste Zeit große Teile der Infrastruktur nicht oder nur wenig ausgelastet sind. Die Untersuchungen von Lanzisera u. a. [La12] zeigen, dass im Jahr 2012 WLAN-Infrastrukturen im Enterprise-Umfeld weltweit mit 2,3 TWh Stromverbrauch zu Buche schlagen. Das macht etwa 2 % des Energieverbrauchs aller Netzwerkkomponenten weltweit aus. Für eine Infrastruktur auf einem Flächengelände wird der anteilige Verbrauch dort mit 5 % angegeben. Gleichzeitig wachsen WLAN-Infrastrukturen um 32 % jährlich ([Ja09]), sodass dem Thema Energieeffizienz auch im WLAN eine nicht unbedeutende Rolle zukommt.

Der naheliegende Ansatz, Teile der unbenutzten Infrastruktur abzuschalten und nur bei Bedarf hochzufahren, steht grundsätzlich im Widerspruch zu der Erwartung der Nutzerinnen und Nutzer, zu jeder Zeit immer vollen und verzögerungsfreien Zugriff auf

das Netzwerk zu haben. Bei der Optimierung des Energieverbrauchs ist daher das Thema der Dienstgüte stets mit im Auge zu behalten.

Ziel dieses Artikels ist es, Strategien für den energieeffizienten Betrieb größerer WLAN-Infrastrukturen auf einem Flächengelände, wie sie bspw. in Bürogebäuden oder auf einem Universitätscampus zu finden sind, bei gleichzeitiger Wahrung der Dienstgüte aufzuzeigen. Am Beispiel der Universität Paderborn wird der IST-Zustand des WLANs bezüglich Last und Energieverbrauch analysiert, um Optimierungspotentiale zu finden und sinnvolle Metriken als Vergleichsgrößen zu identifizieren. Drei unterschiedliche Resource-on-Demand-Algorithmen für den lastabhängigen Betrieb von Access Points werden zur Umsetzung der vorgeschlagenen Strategien vorgestellt und bewertet. Prototypisch werden die Algorithmen auf das WLAN der Universität Paderborn angewendet und der Energieverbrauch sowie die Dienstgüte analysiert. Als Hauptergebnis wird gezeigt, dass die Kombination von unterschiedlichen Betriebsalgorithmen zu unterschiedlichen Zeiten ohne Verlust von Dienstgüte eine Energieeinsparung von 15 % ermöglicht.

Die vorgestellten Strategien, Metriken und Algorithmen sind leicht auf den Betrieb anderer WLAN-Infrastrukturen übertragbar. Sie könnten zusätzlich von Herstellern auch direkt in zukünftige WLAN-Technologien integriert werden.

## 2 Grundlagen

Für das Design und den Betrieb von energieeffizienten WLAN-Infrastrukturen ist es notwendig, sich mit den wichtigsten technischen Eigenschaften der involvierten Komponenten und den physikalischen Zusammenhängen und Effekten auseinanderzusetzen. Dieser Abschnitt gibt einen kurzen Überblick über die wichtigsten Aspekte für die nachfolgenden Kapitel.

Alle erwähnten Informationen betrachten drahtlose Netzwerke nach IEEE 802.11-Standard [Ie12] im Infrastruktur-Modus. Der ebenso spezifizierte Ad-Hoc-Modus verhält sich ähnlich, ist aber nicht Gegenstand dieses Artikels.

Wichtige Akteure in diesen Netzen sind die Komponenten zur zentralen Steuerung und zur Erzeugung einer Infrastruktur, die in der Summe als *Distribution System* (DS) [Re06] bezeichnet werden. Ein DS besteht üblicherweise aus einem oder mehreren *Access Points* (AP), die den Zugangspunkt für Endgeräte (*Mobile Stations* (MS)) bereitstellen. Logische Modellkonstruktionen von 802.11 WLAN-Netzen im Infrastruktur-Modus sind das *Basic Service Set* (BSS), auch als *Zelle* (Cell) bekannt, das alle MS verbunden auf einen AP (inkl. diesem) zusammenfasst, sowie das *Extended Service Set* (ESS), das alle Zellen – und damit auch ihre MS – desselben Netzwerks zusammenfasst. Identifiziert wird ein BSS/ESS durch den *Service Set Identifier* (SSID).

Beim Verbindungsaufbau zu einer WLAN-Infrastruktur sucht die MS anhand einer vorgegebenen SSID nach erreichbaren Zugangspunkten des BSS/ESS und verbindet sich

daraufhin, je nach Implementierung, mit dem Access Point, der mit der höchsten Empfangsqualität (überwiegend SNR, im Folgenden erklärt) empfangen wird.

Die Datenübertragung in drahtlosen Netzwerken nach IEEE 802.11 erfolgt mittels elektromagnetischer Wellen. Dabei interessieren in dieser Arbeit weniger die Verfahren oder die Implementierung, sondern einige Grundlagen in der Signalausbreitung und die Metriken zur Bestimmung der Signalqualität einer Verbindung. Ein einfaches Modell beschreibt Signalstärke als Funktion aus Sendeleistung und Entfernung und geht davon aus, dass sich Sender und Empfänger im freien Raum ohne einschränkende Elemente befinden [Sc05]. In der realen Welt tritt dieser Fall allerdings nahezu nie auf und Funkverbindungen werden u. a. durch Fading, Störungen und Interferenzen gestört [Ku08].

Als Mittel, um Signalqualitäten unter der Beschränkung von Noise und Fading zu beschreiben, kann die Kennzahl *Signal to Noise Ratio* (SNR) benutzt werden. Sollen ebenso Interferenzen berücksichtigt werden, wird von *Signal to Interference and Noise Ratio* (SINR) gesprochen. SNR und SINR können zur Abschätzung der *Bitfehlerrate* (*Bit Error Rate*, BER) benutzt werden, die eine Wahrscheinlichkeit angibt, dass beim Empfänger ein Fehler bei der Extrahierung der empfangenen Daten aus dem Funksignal auftritt [Ku08]. Höhere SNR und SINR sind Anhaltspunkte für niedrigen BER und damit eine zuverlässigere Verbindung. Für gegebene SNR oder SINR hängt die BER ebenso von der Datenrate der Übertragung ab. Höhere Datenraten sind anfälliger für Fehler. Sender-Empfänger-Paare können also auch bei geringen SNR bzw. SINR eine stabile Verbindung durch eine niedrige Datenrate beständig aufrechterhalten [Re06].

Eine im Weiteren verwendete Größe ist die *Received Signal Strength Indication* (RSSI) [Sa10]. Die Kennzahl beschreibt die Verbindungsqualität zwischen Geräten in einem Funknetz. Als Skala wird je nach Quelle eine Ganzzahl zwischen 0 und 255 [Ts06] oder eine logarithmische Skala zwischen -100 dBm und -40 dBm [Ro05] oder bspw. -100 dBm bis -60 dBm [Sa10] bei Mobilfunktechnologien wie UMTS verwendet. Diese Werte sind allerdings mit Vorsicht zu interpretieren, da neben den physikalischen Effekten auch die Charakteristiken der Geräte Einfluss haben. Des Weiteren ist die Menge an Paketverlusten eine mögliche Größe zur Abschätzung der Verbindungsqualität [Ts05]. Hier ist aber auch nur eine Abschätzung der Verbindungsqualität bei bereits schwachen Verbindungen möglich.

Der letzte Aspekt, der für die folgenden Algorithmen erforderlich ist, ist die Messung des Bedarfs der Endgeräte nach Datenübertragungen, um ein Verständnis für die Auslastung der Netze zu entwickeln und auf Engpässe und freie Kapazitäten der Infrastruktur zu reagieren. Eine einfache Kennzahl ist die Anzahl der verbundenen *Mobile Stations* pro *Access Point* [Ma10, Ja07, Ja09]. Hierbei werden allerdings die Charakteristiken der einzelnen Stationen außer Betracht gelassen. So können Access Points mit wenigen bandbreitenintensiven Clients überlastet sein und Access Points mit vielen Geräten, die keine große Bandbreite anfordern, noch Kapazität für weitere Geräte haben. Eine weitere Idee ist das Zählen von TCP-Verbindungen [Ma10]. Diese Strategie hat allerdings auch Nachteile, wenn eine Menge von Clients mit unterschiedlichen Lastprofilen beobachtet werden sollen. Jardosh u. a. [Ja05] definieren *throughput* und



*goodput*. *Throughput* ist die Summe aller Bytes aus allen Frames über eine definierte Zeit. *Goodput* ist die Summe aus allen Bytes aus allen Control-Frames und allen bestätigten Daten-Frames. Eine weitere Metrik, eingeführt in [Ja09], ist die *Channel Utilization*. Sie beschreibt den Prozentsatz der Zeit, in der ein drahtloser Übertragungskanal durch Übermittlungen belegt ist. Hat ein Access Point zwei Kanäle, z. B. im 2,4 GHz- und 5 GHz-Band, ergeben sich unterschiedliche *Channel Utilization*-Werte für beide Kanäle. Überlagern sich Kanäle mit denen anderer Access Points, erhöht sich sehr wahrscheinlich die Kennzahl durch gegenseitige Beeinflussung. Jardosh u. a. [Ja05] weisen darauf hin, dass es einen Zusammenhang zwischen dem Verfahren zur Kollisionsvermeidung (CSMA/CA) in drahtlosen Netzwerken nach 802.11-Standard, der *Channel Utilization*, *throughput* und *goodput* gibt. Bis zu einer Auslastung von 80 % erhöhen sich *throughput* und *goodput*. Eine weitere Erhöhung führt zu einer Vielzahl von Fehlern in den Frames und zu Sendewiederholungen, die die Clients dazu bringen, ihre Datenrate zu drosseln. Durch diese Drosselung sinken *throughput* und *goodput*.

### 3 Analyse der WLAN-Infrastruktur der Universität Paderborn

Um Einsparungspotentiale für den Stromverbrauch von WLAN-Infrastrukturen zu identifizieren, wurde zunächst exemplarisch eine Analyse der vorhandenen Infrastruktur der Universität Paderborn und ihrer Nutzung durchgeführt. Für diese wurden sowohl das Benutzeraufkommen zu verschiedenen Zeiten als auch die vorhandenen Access Points und ihre räumliche Verteilung betrachtet. Die hier beschriebenen Erkenntnisse sind in analoger Form auch auf andere Infrastrukturen übertragbar.

#### 3.1 Benutzeraufkommen

Das Benutzeraufkommen innerhalb einer Infrastruktur ist ein essentieller Aspekt für Versuche, die Energieeffizienz zu erhöhen. Wenn die Sparmaßnahmen die Dienstgüte nicht negativ beeinflussen dürfen, sollten diese nur dann eingesetzt werden, wenn die aktuell verfügbaren Ressourcen von den Benutzern nicht ausgeschöpft werden. Zentrale Systeme wie etwa das in Paderborn verwendete *Aruba AirWave* können einen Überblick über Benutzerzahlen zu verschiedenen Zeitpunkten geben. Abbildung 1 zeigt einen typischen Verlauf der Benutzerzahlen im WLAN der Universität Paderborn über den Zeitraum von einer Woche. Gravierende Unterschiede zwischen Wochentagen und dem Wochenende sowie vom späten Abend bis zu den frühen Morgenstunden sind Indikatoren für mögliche Optimierungsspielräume.

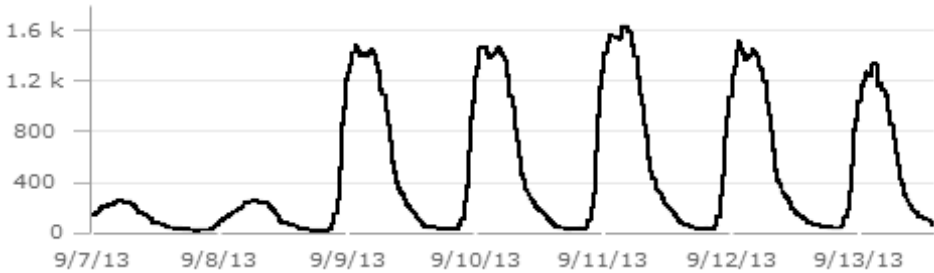


Abbildung 1: Benutzerzahlen wöchentlich (Sa-Fr)

Der gezeigte Verlauf entspricht den Erwartungen für ein universitäres Umfeld. Hierbei sind insbesondere die täglichen Vorlesungszeiten ausschlaggebend. Im Jahresverlauf lassen sich darüber hinaus Unterschiede zwischen normalem Universitätsbetrieb und vorlesungsfreien Zeiten deutlich erkennen. Würde man während der Zeit ohne nennenswerte Nutzung zwischen 23 Uhr und 7 Uhr alle Access Points komplett abschalten, wäre eine Energieeinsparung von 25% realisierbar. Allerdings stände der Dienst dann in dieser Zeit auch nicht zur Verfügung.

### 3.2 Bestandsaufnahme Access Points

Für die Bestandsaufnahme der verwendeten Geräte wurden erneut Daten aus dem *Aruba AirWave*-System verwendet. Insgesamt befanden sich zum Zeitpunkt dieser Arbeit an der Universität Paderborn 393 *Cisco* Access Points diverser Modelle in Betrieb. Die verwendeten Geräte besitzen jeweils zwei Funkmodule und sind so konfiguriert, dass diese jeweils sowohl über das 2.4 GHz- als auch das 5 GHz-Spektrum drahtloses Netzwerk zur Verfügung stellen.

Da die Angaben des Herstellers bezüglich des Stromverbrauchs im Betrieb unter *Power-over-Ethernet* lediglich den Maximalverbrauch enthielten, wurden eigene praktische Messungen durchgeführt, um den *idle*-Verbrauch der Access Points festzustellen. Messungen im Produktiveinsatz von APs konnten keinen wesentlichen Unterschied im Stromverbrauch in Abhängigkeit von aktiven Nutzern oder großer Auslastung feststellen. Der Stromverbrauch aller Access Points an der Universität Paderborn beträgt somit täglich 91.5 kWh.

In der Literatur [Ja07] wurden bei APs an Idle-Zeit Prozentzahlen von 20 bis 60 % zu unterschiedlichen Zeiten und Tagen festgestellt. Die Bestandsaufnahme an der Universität Paderborn hat ergeben, dass durchschnittlich 65 % der vorhandenen APs sich nicht in Nutzung befinden.

### 3.3 Nachbarschaftsbeziehungen

Die räumliche Nähe benachbarter Access Points ist ausschlaggebend für Optimierungsmethoden, welche eine Verringerung der Dienstgüte vermeiden sollen. Die

Distanz zwischen zwei Access Points, unter Beachtung von physikalischen Hindernissen wie z. B. Wänden, ist entscheidend dafür, ob einer dieser Access Points ohne Verlust von Netzabdeckung und ohne zu starke Verringerung der Signalqualität abgeschaltet werden kann. Da ein manuelles Evaluieren hierbei praktisch kaum durchführbar ist, sollten hierfür automatisiert auswertbare Daten verwendet werden.

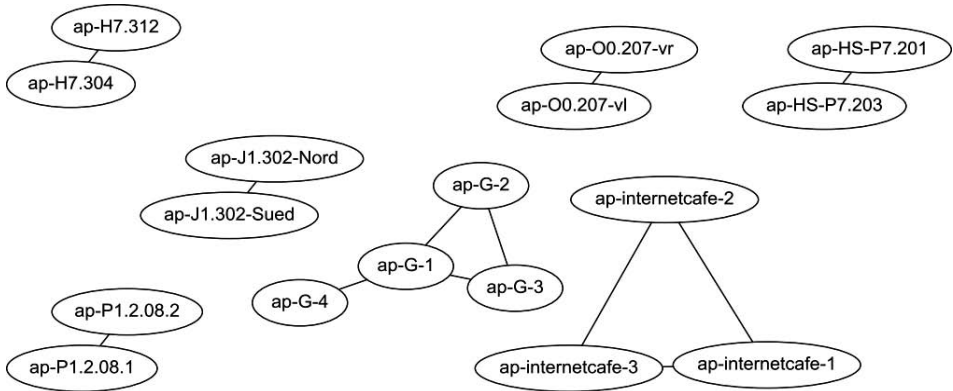


Abbildung 2: Nachbarschaftsbeziehungen (-55 dBm)

Eine Möglichkeit hierfür bieten die RSSI-Signalempfangswerte der Access Points untereinander. In diesem Artikel werden zwei Access Points als Nachbarn angesehen, wenn die Signalempfangswerte beider Funkmodule in beiden Richtungen über einem definierten Schwellenwert liegen. Abbildung 2 zeigt exemplarisch einen Auszug der Nachbarschaftsbeziehungen, die für einen RSSI-Wert von -55 dBm berechnet wurden. Insgesamt ist zu beobachten, dass der Nachbarschafts-Schwellenwert großen Einfluss auf die Anzahl der berechneten Nachbarschaftsbeziehungen hat und demnach ein wichtiger Parameter für etwaige Optimierungsverfahren ist, der sowohl Einfluss auf Einsparungen als auch auf die Dienstgüte hat.

## 4 Optimierungsverfahren

Im Folgenden werden verschiedene Resource-on-Demand-Strategien beschrieben, welche durch gezieltes Abschalten von Access Points Energieeinsparungen erzielen, ohne dabei die Dienstgüte negativ zu beeinflussen.

### 4.1 Green-Clustering

Green-Clustering ist ein von Jardosh u. a. [Ja07] bzw. [Ja09] entwickelter Resource-on-Demand-Algorithmus, welcher benachbarte Access Points zu Clustern zusammenfasst, in denen bedarfsabhängig Access Points an- oder abgeschaltet werden. Als Grundidee dient dabei die Annahme, dass in einem Cluster von benachbarten Access Points bei geringem Bedarf einer von diesen ausreicht, um die gewünschte Netzabdeckung zu gewährleisten. Abbildung 3 zeigt schematisch, wie ein solcher Cluster aus drei

benachbarten Access Points aussehen könnte. Da durch das Abschalten der Nachbarn die Abdeckungsfläche möglicherweise reduziert wird, muss eventuell die Sendeleistung des Cluster-Zentrums erhöht werden, um eine äquivalente Abdeckung zu gewährleisten.

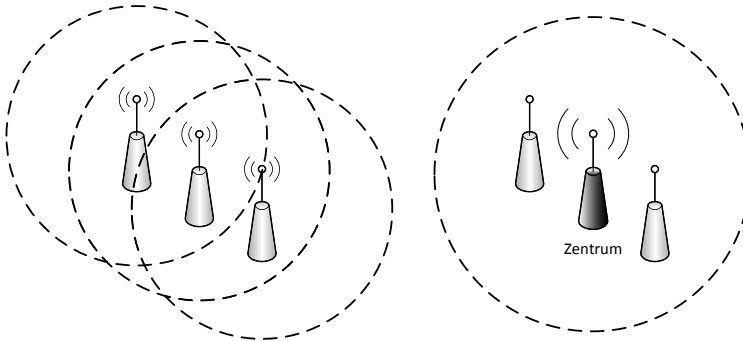


Abbildung 3: Beispielcluster

Abbildung 4 zeigt den Ablauf der Cluster-Bildung nach Jardos [Ja09]. Für die Erstellung der Cluster muss zunächst die sog. *Neighborhood Discovery* durchgeführt werden, bei der alle Nachbarschaftsbeziehungen berechnet werden (siehe [Ja09]). Auf dieser Basis werden anschließend durch einen *Greedy-Algorithmus* Cluster erstellt, wobei die Access Points zunächst nach der Anzahl ihrer Nachbarn sortiert werden und dann für jeden neuen Cluster stets jener Access Point als Cluster-Zentrum ausgewählt wird, welcher die größte Anzahl an Nachbarn besitzt. Anschließend wird der Reihe nach für jeden Nachbarn des Zentrums geprüft, ob dieser auch Nachbar von allen bereits im Cluster befindlichen Access Points ist. Nur wenn dies der Fall ist, wird der betreffende Access Point dem Cluster hinzugefügt. Ist ein Cluster komplett, werden die betreffenden Access Points nicht länger als Nachbarn anderer Access Points angesehen. Anschließend werden die verbleibenden Access Points erneut nach der Anzahl ihrer Nachbarn sortiert und ein weiterer Cluster erstellt. Der Algorithmus terminiert, wenn keine weiteren Access Points ohne Cluster verbleiben, welche noch Nachbarn besitzen.

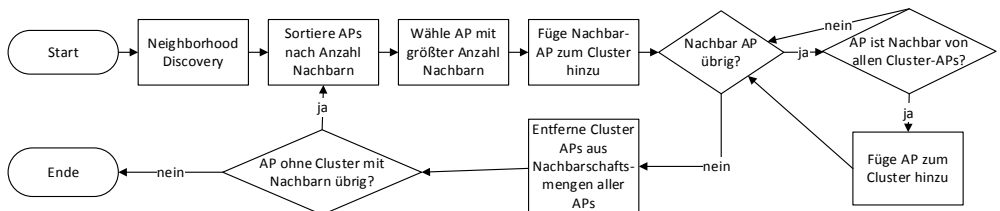


Abbildung 4: Green-Cluster Greedy-Algorithmus

Im Anschluss an das Clustering können die Cluster im laufenden Betrieb verwendet werden, um die Infrastruktur dynamisch an den Bedarf anzupassen. Dabei bleiben Cluster-Zentren und Access Points, welche keinem Cluster zugeordnet werden konnten, stets angeschaltet. Innerhalb der Cluster verwenden Jardosh u. a. [Ja09] die gemessene *Channel-Utilization* als Metrik zur Bedarfsermittlung. Überschreitet dieser Wert für einen Access Point innerhalb eines Clusters einen vorgegebenen Schwellenwert, wird, sofern möglich, ein Cluster-Nachbar zugeschaltet. Falls ein zugeschalteter Cluster-AP

für eine bestimmte Zeit keine verbundenen Benutzer hat, wird dieser abgeschaltet. Um die Last innerhalb eines Clusters zu verteilen, werden die verbundenen Benutzer außerdem durch ein aktives *Load Balancing* auf die Cluster-APs verteilt.

Der *Greedy*-Algorithmus von Jardosh u. a. wurde für die Infrastruktur der Universität Paderborn mit einigen kleinen Änderungen implementiert und evaluiert. Die Ergebnisse der Implementierung waren positiv, zeigten jedoch auch Verbesserungspotential für die Cluster-Bildung auf. Der nächste Abschnitt beschreibt, wie das Verfahren unter Verwendung eines mathematischen Optimierungsmodells verbessert wurde.

## 4.2 IP-Green-Clustering

Da der *Greedy*-Algorithmus nachweislich (siehe [Bo13]) nicht-optimale Cluster liefern kann, wurde ein ganzzahliges mathematisches Optimierungsmodell zur Cluster-Bildung formuliert und unter Benutzung der *Java*-Schnittstelle des *GLPK*-Solvers implementiert. Die Modellformulierung sieht wie folgt aus:

Mengen und Parameter:

---

|                   |                                                      |
|-------------------|------------------------------------------------------|
| $A$               | Menge der Access Points                              |
| $n_{ij} \in 0, 1$ | Parameter, der angibt, ob $i, j \in A$ Nachbarn sind |

Entscheidungsvariablen:

---

|          |                                                                                                     |
|----------|-----------------------------------------------------------------------------------------------------|
| $y_{ij}$ | Binär-ganzzahlige Variable, die angibt, ob $i, j \in A \wedge i \neq j$ demselben Cluster angehören |
|----------|-----------------------------------------------------------------------------------------------------|

Zielfunktion:  $\max z = \sum_{i \in A} \sum_{j \in A} y_{ij}$

Restriktionen:

- 1)  $y_{ij} \leq n_{ij} \forall i, j \in A \wedge i \neq j$
- 2)  $y_{ij} - y_{ji} \leq 0 \forall i, j \in A \wedge i \neq j$
- 3)  $y_{ij} - y_{jk} \leq 1 \forall i, j, k \in A : n_{ik} = 0 \wedge i \neq j \neq k$
- 4)  $y_{ij} + y_{ik} - y_{jk} \leq 1 \forall i, j, k \in A : n_{ik} = 1 \wedge i \neq j \neq k$
- 5)  $y_{ij} \in \{0, 1\} \forall i, j \in A, i \neq j$

Für jedes mögliche Paar von Access Points existiert eine Entscheidungsvariable, welche angibt, ob die beiden APs demselben Cluster angehören oder nicht. Durch das Maximieren dieser Variable in der Zielfunktion wird die Anzahl der „aktivierten“ Cluster-Verbindungen maximiert. Die Restriktionen 2), 3) und 4) stellen hierbei sicher, dass jeder Access Point maximal einem Cluster angehören kann und innerhalb der Cluster alle APs Nachbarn sind.

Für die betrachteten Instanzen liefert das Modell innerhalb einer Laufzeit von wenigen Sekunden optimale Lösungen bezüglich der für die Grundabdeckung benötigten APs. Theoretisch existieren jedoch auch Fälle, für welche das Modell nicht-optimale Lösungen liefert. Dieses Problem könnte durch eine komplexere Modellformulierung gelöst werden, welche jedoch wahrscheinlich zu längeren Laufzeiten führen würde. Die Implementierung liefert i. d. R. leicht verbesserte Ergebnisse im Vergleich zum ursprünglichen Green-Clustering. Der nachfolgende Abschnitt beschreibt ein Verfahren, welches weitere Einsparungen ermöglicht.

### 4.3 Green-Star-Clustering

Das Ergebnis der bereits beschriebenen Verfahren ist, dass innerhalb einer Infrastruktur jeder Access Point entweder angeschaltet ist oder einen Nachbarn hat, der aktuell dessen Abdeckungsbereich übernimmt. Konstellationen wie in Abbildung 5 ermöglichen dem Green-Clustering kaum Einsparungen, obwohl es offensichtlich möglich ist, durch jeweils nur einen angeschalteten AP die nötige Grundabdeckung zu erzielen. Wie in der Abbildung zu sehen ist, würde es genügen, die beiden zentral gelegenen APs anzuschalten, um die sternförmig darum angeordneten Nachbarn mitzuversorgen.

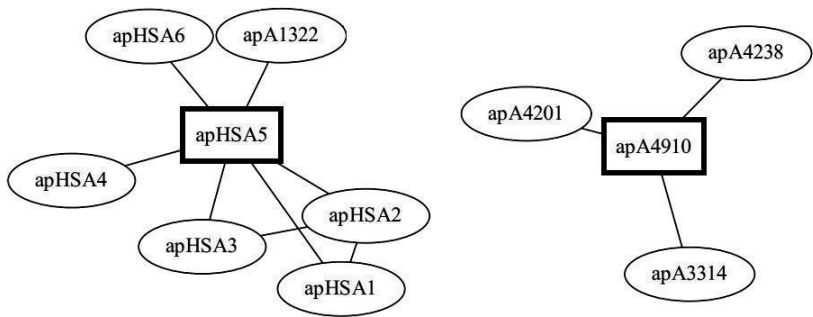


Abbildung 5:Green-Star-Clustering

Das komplett neu entwickelte Green-Star-Clustering ermöglicht es, die Anzahl der benötigten Access Points weiter zu reduzieren. Auch dieses Verfahren wurde unter Verwendung eines ganzzahligen Optimierungsmodells implementiert.

Das Modell minimiert die Anzahl der permanent angeschalteten APs. Gleichzeitig sorgt Restriktion 1) dafür, dass jeder Access Point entweder selber angeschaltet ist oder einen angeschalteten Nachbarn hat.

Mengen und Parameter:

---

$A$  Menge der Access Points

$N_i$  Menge der benachbarten APs  $j \in A$  des APs  $i \in A$

Entscheidungsvariablen:

---

$y_i$  Binär-ganzzahlige Variable, die angibt, ob  $i \in A$  permanent angeschaltet ist oder nicht

Zielfunktion:  $\min z = \sum_{i \in A} y_i$

Restriktionen:

- 1)  $y_i + \sum_{j \in N_i} y_j \geq 1 \quad \forall i \in A$
- 2)  $y_i \in \{0,1\} \quad \forall i \in A$

Verglichen mit Green-Clustering und IP-Green-Clustering liefert dieses Modell stark verbesserte Ergebnisse für die minimal benötigten permanent angeschalteten APs. Allerdings fehlen bei diesem Verfahren „echte“ Cluster, in denen jeder AP die Aufgaben der anderen übernehmen kann. Im laufenden Betrieb führt dies zu einer erhöhten Komplexität, da bei einer Bedarfssteigerung anstatt nur eines APs eventuell mehrere hinzugeschaltet werden müssen. Da die evaluierten Access Points mit 2-3 Minuten zum Hochfahren auch verhältnismäßig lange brauchen und Antennen nicht einzeln abgeschaltet werden können, kann es daher zum Verlust der Dienstgüte kommen.

Die einzelnen Algorithmen und die Messung der Dienstgüte werden im Detail in [Bo13] beschrieben.

## 5 Implementierung

Die oben beschriebenen Algorithmen wurden in Java implementiert und mit realen Daten der Infrastruktur der Universität Paderborn getestet. Hierfür wurden zunächst per *SNMP*-Schnittstelle aktuelle AP-Statusinformationen von den WLAN-Controllern extrahiert. Auf Basis dieser Daten wurden danach Cluster berechnet. Anhand von Signaldaten wurde anschließend analysiert, wie sich das Abschalten von Access Points auf die Signalqualität der Client-Geräte auswirken würde. Das An- bzw. Abschalten von Access Points wurde über einen Web-Service realisiert, welcher es ermöglicht, Ethernet Ports an den jeweiligen Switches einzeln zu schalten.

## 6 Zusammenfassung

In der Praxis empfiehlt es sich nach Analyse der Ergebnisse, den Green-Star-Clustering verstärkt zu Zeiten einzusetzen, wo typischerweise wenig Benutzer- bzw. Bedarfsfluktuation herrscht. Auf Basis einer detaillierten Analyse der Nutzungsaufkommen zu verschiedenen Tageszeiten sollte die WLAN-Infrastruktur an der Universität Paderborn demnach tagsüber von 7 bis 21 Uhr per IP-Green-Clustering und über Nacht von 21 bis 7 Uhr des nächsten Tages vom Green-Star-Clustering verwaltet werden. Damit ergibt sich eine Energieeinsparung von 15 % bei gleichbleibender Dienstgüte im Vergleich zu einer Einsparung von 25 % bei einer Komplettabschaltung zwischen 23 und 7 Uhr.

Die Ergebnisse wurden im Rahmen des GreenPAD-Projekts (gefördert vom BMWi im Rahmen der IT2Green-Ausschreibung) erzielt.

## Literaturverzeichnis

- [Bo13] Boehner, M.: Optimization the Energy-Consumption of WLAN-Infrastructures, Master thesis, Universität Paderborn, 2013
- [Ie12] IEEE: IEEE Std 802.11 - 2012: Telecommunications and information exchange between systems - local and metropolitan area networks - specific requirements, 2012
- [Ja05] Jardosh et al.: Understanding congestion in IEEE 802.11b wireless networks. IMC '05 Proceedings of the 5th ACM SIGCOMM conference on Internet Measurement, 2005
- [Ja07] Jardosh et al.: Towards an Energy-Star WLAN Infrastructure. Mobile Computing Systems and Applications, 2007
- [Ja09] Jardosh et al.: Green WLANs: On-Demand WLAN Infrastructures. Mobile Networks and Applications, 2009
- [Ku08] Kumar, et al.: Wireless networking. Morgan Kaufmann/Elsevier, 2008
- [La12] Lanzisera, S.; Nordman, B.; Brown, R.E.: Data network equipment energy use and savings potential in buildings. *Energy Efficiency*, 5(2), 2012, S. 149-162
- [Ma10] Marsan et al.: A simple analytical model for the energy-efficient activation of access points in dense WLANs. Proceedings of the 1st International Conference on Energy-Efficient Computing and Networking, 2010
- [Re06] Rech, Jörg: Wireless LANs: 802.11-WLAN-Technologie und praktische Umsetzung im Detail; 802.11a/h, 802.11b, 802.11g, 802.11i, 802.11n, 802.11d, 802.11e, 802.11f, 802.11s. Heise, 2006
- [Ro05] Rodrig et al.: Measurement-based characterization of 802.11 in a hotspot setting. Proceedings of the 2005 ACM SIGCOMM workshop on Experimental approaches to wireless network design and analysis, 2005
- [Sa10] Sauter, M.: From GSM to LTE: An Introduction to Mobile Networks and Mobile Broadband. Wiley, 2010
- [Sc05] Schwartz, Mischa: Mobile wireless communications. Cambridge University Press, 2005
- [Ts05] Tsukamoto et al.: Impact of layer 2 behavior on TCP performance in WLAN. IEEE VEHICULAR TECHNOLOGY CONFERENCE, 2005
- [Ts06] Tsukamoto et al.: Experimental Evaluation of Decision Criteria for WLAN Handover: Signal Strength and Frame Retransmission. Ubiquitous Computing Systems Journal, 2006





# AEQUO - Adaptive und energieeffiziente Verteilung von virtuellen Maschinen in OpenStack-Umgebungen

Kai Spindler, Sebastian Rieger

Fachbereich Angewandte Informatik

Hochschule Fulda

Marquardtstraße 35

36039 Fulda

[vorname.nachname]@informatik.hs-fulda.de

**Abstract:** Servervirtualisierungs- und Cloud-Infrastrukturen nehmen einen immer größeren Teil der IT-Infrastruktur von Rechenzentren ein. Für Rechenzentren sind dabei insbesondere die laufenden Betriebskosten (vorrangig für Strom und Klimatisierung) ein zunehmender Faktor. Lösungen für eine Steigerung der Energieeffizienz in Cloud-Infrastrukturen ermöglichen eine Reduzierung dieser Betriebskosten und eine optimale Auslastung der Rechenzentrumsinfrastruktur. In der vorliegenden Arbeit wird eine Lösung beschrieben, die Leistungsmetriken und insbesondere Energieeffizienzparameter in einer OpenStack-Cloud-Infrastruktur überwacht und durch effiziente Verlagerung von virtuellen Maschinen optimiert. Bestehende OpenStack-Infrastrukturen können durch die entwickelte Software-Komponente um eine adaptive Verteilung der darin betriebenen virtuellen Maschinen erweitert werden.

## 1 Einleitung

In den vergangenen Jahren hat die Verwendung von Cloud-Diensten insbesondere durch die anwachsende Zahl mobiler Endgeräte (Smart Phones, Tablets) mit denen die Dienste genutzt werden weiter zugenommen. Um einen schnellen und fehlerfreien Zugriff auf die Dienste zu gewährleisten, sind dabei im Hintergrund leistungsstarke Rechenzentren erforderlich. Dies bedeutet mit Blick auf die Kosten für Strom und Klimatisierung in Rechenzentren eine immer größere Herausforderung. Hierfür sind insbesondere leistungsfähige IT-Virtualisierungslösungen und Rechnernetze erforderlich, um die Dienste effizient und leistungsstark über das Internet zur Verfügung zu stellen. In Bezug auf die Effizienz sind hierbei vor allem die Kosten für Strom und Klimatisierung der erforderlichen IT-Ressourcen (z.B. Server-, Storage-, Netzkomponenten) zu betrachten. Ein lastabhängiges Zu- und Abschalten dieser Ressourcen bzw. eine adaptive Steuerung von deren Leistungsaufnahme (vgl. Drosselung der Taktgeschwindigkeit, Energiesparfunktionen etc.) kann somit helfen die Kosten für den Betrieb von virtualisierten IT-Infrastrukturen bzw. virtuellen Maschinen zu reduzieren und eine gleichmäßige Auslastung zu erzielen. Darüber hinaus können adaptive Verfahren, die die virtuellen Maschinen und Cloud-Dienste je nach Anforderung skalieren, auch deren Leistungsfähigkeit insgesamt optimieren.

Dieses Paper präsentiert einen entsprechenden Ansatz, um in einer OpenStack-basierten Cloud-Infrastruktur virtuelle Maschinen im Hinblick auf die Energieeffizienz zu verteilen und zu skalieren. Hierfür wurde ein Prototyp entwickelt, der zum einen eine Monitoring-Funktion (als Sensor) implementiert, bzw. für die Energieeffizienz relevante Parameter innerhalb der OpenStack-Infrastruktur (Temperatur resp. Kühlung, Auslastung von Compute, Storage und Netzkomponenten) überwacht, und zum anderen eine Management-Funktion (als Aktor) realisiert, die virtuelle Maschinen über die einzelnen Nodes der Cloud-Infrastruktur verteilen, verlagern sowie skalieren kann. Hierbei können zusätzlich einzelne Funktionen bzw. Komponenten innerhalb der Infrastruktur anhand des aktuellen Energiebedarfs zu- oder abgeschaltet bzw. gedrosselt werden. Der Prototyp verwendet für die Steuerung der Cloud-Ressourcen die Standard OpenStack API. Zusätzlich wird die Nutzung des Open Cloud Computing Interface (OCCI) Standards evaluiert, wodurch eine Adaptierung für andere Cloud-Umgebungen bzw. eine Anwendung in Cloud-Föderationen bestehend aus unterschiedlichen Cloud-Anbietern möglich wird.

Die folgenden Abschnitte geben eine Einführung in verteilte OpenStack-Umgebungen und die Energieeffizienz von Private Cloud-Umgebungen. Abschnitt 2 geht auf Anforderungen an eine energieeffiziente Verteilung von virtuellen Maschinen in OpenStack-Umgebungen ein. Dabei werden Anknüpfungspunkte an die standardmäßig verfügbare Verteilung in OpenStack sowie der aktuelle Stand der Forschung in diesem Bereich beschrieben. Anschließend beschreibt Abschnitt 3 die von unserem Prototyp überwachten Energieeffizienzparameter in OpenStack-Umgebungen. Dies umfasst die Umsetzung der in diesem Paper genannten Anforderungen sowie erste Auswertungen des Testbetriebs. Abschließend runden ein Fazit sowie ein Ausblick dieses Paper ab.

## **1.1 Virtualisierung in verteilten OpenStack-Umgebungen**

OpenStack [Ope14] bietet eine Open Source Lösung für die Realisierung von Public und Private Cloud-Umgebungen. Ein weit verbreitetes Klassifizierungsmodell für Cloud Computing bildet die Definition des NIST [MG11]. Hierbei werden fünf essenzielle Charakteristika, drei Service-Modelle und vier Bereitstellungsmodelle des Cloud Computing unterschieden. Im Fokus dieses Papers stehen in Bezug auf die Bereitstellungsmodelle Private Clouds, da hier durch den Betrieb im eigenen Rechenzentrum die Energieeffizienz bzw. Energiekosten eine zentrale Rolle spielen. OpenStack adressiert primär „Infrastructure as a Service“ (IaaS), so dass dieses Service-Modell im Vordergrund dieses Papers steht. In Bezug auf die Charakteristika beschreibt das NIST Paper die Cloud als einen einfachen, ubiquitären, on-demand über das Netz bereitgestellten Pool aus konfigurierbaren Computing Ressourcen: Server, Storage, Netze, Anwendungen und Dienste. Die Infrastruktur bzw. IaaS bilden hiervon nur Server (OpenStack Nova), Storage (OpenStack Swift und Cinder) und Netze (OpenStack Neutron) ab, deren Energieverbrauch somit in diesem Paper optimiert wird. Als Anwendungen bzw. Dienste wird darauf aufbauend vorrangig der Betrieb von virtuellen Maschinen (VMs) als Nova Instances auf den Servern der OpenStack-Umgebung betrachtet. VMs können hierbei auf bestimmten Servern platziert werden sowie innerhalb ihres Lebenszyklus auf andere Server migriert oder gestoppt wer-

den. Auch eine Skalierung der Leistungsfähigkeit der VMs wird unterstützt. VMs können dabei auf Server in verschiedenen Racks oder auch über verschiedene Standorte verteilt und einheitlich innerhalb der OpenStack-Umgebung verwaltet werden.

## 1.2 Energieeffizienz von Private Cloud-Umgebungen

Die Energieeffizienz von Rechenzentren und IT-Systemen ist nicht zuletzt seit der Einführung des Schlagworts „Green-IT“ vermehrt in den Vordergrund gerückt [HS10] [HF12]. Ein wesentliches Problem bildet die Verlustleistung der betriebenen Systeme (vgl. Power Usage Effectiveness, kurz PUE [WDD12]) bzw. die Effizienz, mit der der Strom tatsächlich für die vom Rechenzentrum betriebenen Dienste und nicht für die zusätzlich notwendige Infrastruktur darum herum verwendet wird. Darüber hinaus stellen steigende Energiekosten sowie Folgekosten für die Klimatisierung der Systeme und unterbrechungsfreie Stromversorgung eine große Herausforderung für den Betrieb moderner und immer größer werdender Rechenzentren dar. Der Energiebedarf innerhalb von OpenStack-Umgebungen wird in Bezug auf diese Kriterien insbesondere hardwareseitig durch die im vorherigen Abschnitt genannten IaaS Komponenten Server (z.B. CPU, RAM, NIC), Storage (z.B. HDD, NAS, SAN) und Netze (z.B. Router, Switches, Links) bestimmt. Darüber hinaus ist die Effizienz der innerhalb dieser Komponenten verwendeten Spannungswandler bzw. Netzteile ein wesentlicher Faktor, sowie die erforderliche Energie für die Abführung der Verlustleistung resp. Klimatisierung. Die detaillierte Messung der Energieeffizienz beim Betrieb von virtuellen Maschinen wurde bereits von einigen verwandten Arbeiten wie z.B. [VT13] bewertet. Weitere verwandte Arbeiten in diesem Bereich werden im Abschnitt 2.3 genannt. In Bezug auf die fünf charakteristischen Eigenschaften von Cloud-Lösungen nach [MG11], wird die Energieeffizienz von Private Cloud-Umgebungen in diesem Paper definiert als:

**On-demand self-service** Energie muss für die eigenständig von den Kunden verwalteten Systeme unmittelbar und nach Bedarf zur Verfügung gestellt werden, ohne dass Anpassungen durch das Rechenzentrum erforderlich werden.

**Broad network access** Cloud-Services werden dezentral über Netze, insb. dem Internet, durch eine Vielzahl von Anwendern genutzt. Hohe Zugriffsbandbreiten treiben hierbei auch die Energiekosten (z.B. pro aktivem Port eines Switches inkl. SFP, Linkstrecke etc.) in die Höhe. In Bezug auf die Netzinfrastruktur können entfernte VMs über eine Vielzahl von Vermittlern (Router, Switches) topologisch miteinander gekoppelt sein, wodurch der Energiebedarf einer VM indirekt steigt.

**Resource pooling** Komponenten werden gemeinsam von mehreren Kunden verwendet. Dies ermöglicht eine Aufteilung und Einsparung von Energiekosten.

**Rapid elasticity** Bedingt durch das Resource Pooling muss u.U. leistungsfähigere Hardware verwendet werden, die nur dann in Bezug auf die Anschaffungs- und Betriebskosten (inkl. Strom, Klimatisierung) effizient ist, wenn sie optimal ausgelastet wird.

**Measured service** Energiekosten sollten gegenüber dem Kunden transparent und feingranular abgerechnet werden können.

OpenStack Server (Compute Nodes) greifen auf unterschiedliche Storage- und Netz-Komponenten zu. Abhängig von der Position, Ausführung und Konfiguration dieser Komponenten entsteht ein unterschiedlicher Energieverbrauch. Dies erhöht insb. den Overhead und damit die Energiekosten für Redundanzlösungen wie z.B. zusätzliche Rechenzentrumsstandorte, deren Anbindung (Data Center Interconnect) und den Abgleich der Services und Daten über die verteilten Lokationen (Replikation etc.). Hierbei können auch schwankende Energiekosten an den Standorten (vgl. Verfügbarkeit von erneuerbaren Energien oder Smart Grid) sowie unterschiedliche klimatische Bedingungen (bzw. damit verbundene Kosten für die Klimatisierung) Einfluss auf die energieeffiziente Platzierung und Verlagerung von virtuellen Maschinen in Private Cloud-Umgebungen haben.

## 2 Anforderungen an die energieeffiziente Verteilung von VMs

Bezüglich der energieeffizienten Verteilung von virtuellen Maschinen in OpenStack-Umgebungen werden in den folgenden Abschnitten die standardmäßig von OpenStack unterstützten Verteilungskriterien erläutert sowie Anforderungen für eine Steigerung der Energieeffizienz und Ergebnisse verwandter Arbeiten in diesem Bereich diskutiert.

### 2.1 Platzierung von virtuellen Maschinen in OpenStack durch den Nova Scheduler

OpenStack Nova bietet standardmäßig mehrere Verfahren für die Platzierung von virtuellen Maschinen auf den Compute Nodes. Diese werden durch den *nova-scheduler* [Fou14b] definiert. Default-Einstellung ist der sog. *FilterScheduler*, der Compute Nodes für die VMs anhand von deren Standort, verfügbarem RAM und unterstützten Funktionen und Ressourcen (z.B. bestimmte Storage- oder Netzwerk-Anforderungen) auswählt. Zusätzlich können weitere Filterkriterien, wie z.B. die Trennung von VMs auf zwei unterschiedlichen physikalischen Hosts oder umgekehrt deren Gruppierung auf einem gemeinsamen Host definiert werden. Auch die Migration von VMs kann durch den Scheduler unterstützt werden. Sowohl bei der Platzierung als auch bei der Migration werden allerdings von Nova keine Energieeffizienzparameter in die Scheduling-Entscheidung einbezogen. Auch eine automatische Optimierung bzw. Ausbalancierung der Auslastung durch entsprechende Migrationen zur Steigerung der Energieeffizienz wird standardmäßig nicht unterstützt.

### 2.2 Optimierte Platzierung von VMs im Hinblick auf die Energieeffizienz

Naheliegende Metriken in Bezug auf die effiziente Verteilung von virtuellen Maschinen (VMs) in einer Virtualisierungsumgebung sind die CPU- sowie RAM-Auslastung der Hy-

pervisoren. Ein einfaches Beispiel für die Veranschaulichung von deren Relevanz bilden zwei Hypervisoren auf denen nur eine VM läuft und die Hypervisoren jeweils zu 10% auslastet. Es ist offensichtlich, dass in diesem Szenario unnötig Energie verschwendet wird, da beide VMs auch gemeinsam auf einem einzelnen Hypervisor laufen könnten. Der andere Hypervisor könnte dann in einen Ruhezustand (vgl. ACPI S3 oder S4) versetzt werden. In der Realität sind entsprechende Szenarien oft weitaus komplexer. So befinden sich in der Regel mehrere VMs auf einem Hypervisor und es werden Auslastungen von ca. 80% oder mehr angestrebt. Häufig werden jedoch auch obere Grenzen für die Konsolidierung und somit Auslastung definiert, um Kapazitätsengpässe bei Lastspitzen zu vermeiden. Um Kapazitätsengpässe zu vermeiden, ist es zusätzlich notwendig Vorhersagen über die Auslastung zu treffen. Hierfür wurden bereits Algorithmen in verwandten Arbeiten präsentiert, die in Abschnitt 2.3 genannt werden. Der in diesem Paper präsentierte Prototyp ermöglicht die Überwachung von Metriken bzgl. des aktuellen Energieverbrauchs in Cloud-Umgebungen, die dann durch diese Algorithmen verwendet werden können, um eine optimale und energieeffiziente Platzierung der VMs zu unterstützen.

Die Abwärme von Servern ist ein wesentlicher Kostenfaktor beim Betrieb von Rechenzentren, da ein Hauptteil der verbrauchten Energie eines Rechenzentrums für deren Kühlung verwendet wird. Weiterhin können durch die immer engere Bauweise von Servern und die steigende Leistungsaufnahme immer leistungsfähigerer Hardware Probleme entstehen, bei denen die Hitze nicht mehr ausreichend abgeführt werden kann. Dies kann zur Verkürzung der Lebensdauer der Systeme beitragen. Eine gute Verteilung des Kühlaufwands in einem Rechenzentrum, z.B. unterstützt durch eine optimale Verteilung von VMs, kann zur Senkung der Kosten für die Kühlung beitragen [Bel13].

Durch eine Überwachung des Netzwerkverkehrs wird darüber hinaus ermöglicht die Energiekosten zu senken, indem VMs die oft miteinander kommunizieren, auf dem gleichen bzw. auf nahe beieinander liegenden Compute Nodes platziert werden. Außerdem können durch eine intelligente Verteilung Teile des Netzwerks abgeschaltet (z.B. Ports, Links) oder in ihrer Leistung gedrosselt werden. Standards zu Energy Efficient Ethernet (802.3az) sind beispielsweise in der Lage Ports mit weniger Traffic in einen Modus zu versetzen, bei dem auf Kosten höherer Latenz weniger Strom verbraucht wird. Durch ein adaptives Management dieser Funktion lässt sich die Energieeffizienz steigern ohne dauerhaft höhere Latenzen in Kauf nehmen zu müssen. In Umgebungen, die ein Multipathing einsetzen, können zudem in Zeiten niedriger Netzauslastung Netzwerkkomponenten (Switches, Router, Links) abgeschaltet bzw. zusätzliche Energiesparfunktionen genutzt werden.

Erneuerbare Energien können darüber hinaus in die energieeffiziente Verteilung von VMs in Cloud-Umgebungen einbezogen werden. Beispielsweise könnten so günstige Energiepreise durch einen Überschuss von Wind- oder Solarenergie an einzelnen Standorten genutzt werden. Zusätzlich könnten Temperaturunterschiede an den Standorten durch eine entsprechende Verschiebung der VMs die Klimatisierungs- und damit Energiekosten weiter reduzieren. Möglich wird diese transparente Verlagerung der VMs durch zunehmend verfügbare standortübergreifende Layer 2 Netztopologien sowie neue Ansätze für deren adaptives Management. Entsprechende Lösungen wie z.B. Software-defined networking können die Netzwerklast anwendungsabhängig auf Flow-Ebene steuern und damit auch eine Steigerung der Energieeffizienz erzielen [HSM<sup>+</sup>10].

## 2.3 Verwandte Arbeiten

Die energieeffiziente Verteilung von virtuellen Maschinen innerhalb von OpenStack-Umgebungen wird ebenfalls in [Bel13] und [BB10] betrachtet. Allerdings wird in diesen Ansätzen keine temperatur- oder netzabhängige Platzierung der VMs umgesetzt. Zusätzlich lässt sich die dort entwickelte Erweiterung nicht im aktuellen Havana Release von OpenStack verwenden. Darüber hinaus ist eine Integration weiterer Entscheidungskriterien in den Scheduling-Prozess bzw. die Verwendung von Standard-APIs und Komponenten von OpenStack nicht vorgesehen. Eine allgemeinere und detaillierte Bewertung der für die Energieeffizienz von Cloud-Umgebungen relevanten Faktoren findet sich in [SFW<sup>+</sup>13] und [SH13]. Diese bieten allerdings keine Testumgebung für OpenStack. Faktoren für den konkreten Energiebedarf von VMs sowie deren Migration werden z.B. in [VT13] präsentiert.

Die Optimierung der Energieeffizienz von Rechenzentrumsnetzen durch eine adaptive Platzierung von VMs wurde bereits in [MKDK11], [FLL<sup>+</sup>13] und [AG13] diskutiert. In diesem Paper wird ein Prototyp für OpenStack vorgestellt, der diese Ansätze mit der Überwachung und Steuerung von Power Management Funktionen wie z.B. ACPI und 802.3az kombiniert, wie sie z.B. in [CRN<sup>+</sup>10] evaluiert wurden. Der Prototyp ermöglicht durch die Kombination bestehender Power Management Lösungen und eine adaptive Zu- bzw. Abschaltung und Drosselung der Cloud-Ressourcen eine Steigerung der Energieeffizienz in OpenStack-basierten Private Cloud-Umgebungen. Dabei wird neben dem Power Management auch eine Überwachung der Temperatur in Rechenzentren unterstützt, um so zusätzlich die Klimatisierung zu optimieren und die Energiekosten weiter zu senken.

## 3 Adaptive Verteilung von VMs in OpenStack durch AEQUO

AEQUO (lat. ausgewogen) bezeichnet einen Prototypen innerhalb eines Forschungsprojekts der Hochschule Fulda, der u.a. das energieeffiziente Verlagern von virtuellen Maschinen in OpenStack unterstützt. Im Folgenden werden die durch den Prototypen umgesetzten Anforderungen sowie die realisierte Testumgebung erläutert.

### 3.1 Testumgebung

Für das Deployment von OpenStack erwies sich die Rackspace Private Cloud [Clo14] als tragfähige Lösung. Im nachfolgend beschriebenen Proof of Concept wurden drei Rechner verwendet, wovon zwei als dedizierte Nova Compute Nodes dienen. Der dritte Rechner wurde als Hypervisor (VMware ESXi) installiert, der virtuelle Maschinen (Ubuntu 12.04 LTS) für den Nova Controller, den Chef Server (zuständig für das Deployment) sowie Cinder (Block Storage) und Neutron (Networking) bereitstellt. Abbildung 1 zeigt die in der Testumgebung verwendeten Funktionskomponenten.

Um zwischen den Nova Compute Nodes die VMs während des Betriebs unterstützt durch

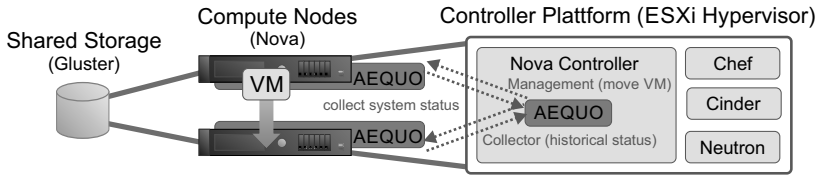


Abbildung 1: Architektur der AEQUO Testumgebung

die AEQUO Komponenten umziehen zu können (Live-Migration) ist es notwendig, einen Shared Storage zwischen den Nova Compute Nodes einzurichten [Fou14a]. Um den Shared Storage für den Nova Controller und die Compute Nodes zur Verfügung zu stellen wurde Gluster [GLU14] verwendet.

### 3.2 Realisierung

Für den Proof of Concept wurde AEQUO als Python Script entwickelt. AEQUO besteht aktuell aus drei Bestandteilen. Der erste Teil ist der *Collector Daemon*, welcher auf jedem Nova Compute Node läuft und dort Leistungsmetriken (z.B. Temperatur, Last) sammelt. Die anderen zwei Komponenten werden auf dem Nova Controller ausgeführt, wobei der *Aggregator Daemon* die Leistungsdaten der Controller sammelt und in eine SQLite Datenbank schreibt. Die dritte Komponente, der *Balancing Daemon*, läuft ebenfalls auf dem Nova Controller, holt sich regelmäßige Leistungsdaten der VMs zu den letzten x Messpunkten aus der Datenbank und sorgt gegebenenfalls dafür, dass VMs umgezogen werden. Abbildung 2 verdeutlicht diesen Aufbau.

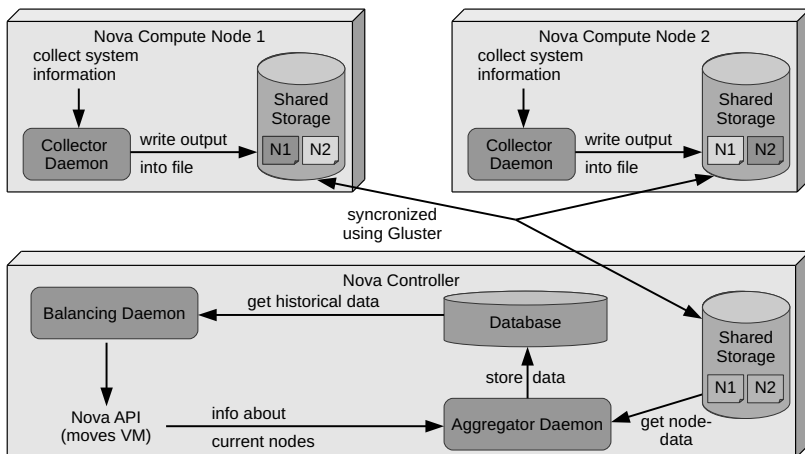


Abbildung 2: AEQUO Aufbau und Architektur

Die Datenbank in die der *Aggregator Daemon* die Daten schreibt, besteht hauptsächlich aus drei Tabellen, siehe Abbildung 3. Die Tabelle „measure“ beinhaltet die historischen Messdaten, welche später genutzt werden, um Entscheidungen zu treffen ob eine VM um-



gezogen wird oder nicht. Die Tabelle „nodedata“ enthält generelle Informationen zu den einzelnen VMs. In die dritte Tabelle „vmmove“ wird ein Log über die durchgeführten Umzüge geschrieben.

| Feld     | Typ  | Feld     | Typ     | Feld       | Typ     |
|----------|------|----------|---------|------------|---------|
| hostname | TEXT | hostname | TEXT    | hostname   | TEXT    |
| ip       | TEXT | time     | INTEGER | time       | INTEGER |
| status   | TEXT | temp1    | REAL    | moved_from | TEXT    |
|          |      | ...      |         | moved_to   | TEXT    |

(a) Tabelle: nodedata                      (b) Tabelle: measure                      (c) Tabelle: vmmove

Abbildung 3: Datenbankstruktur von AEQUO

Der *Collector Daemon*, welcher auf den Compute Nodes installiert wird, benötigt keine gesonderte Konfiguration, da er die Leistungsmetriken direkt aus dem System sammelt und anschließend in eine Datei im Shared Storage schreibt. Die gesammelten Metriken werden im JSON Format in den Dateien abgelegt. Hierbei legt jeder Compute Node seine eigene Datei an. Für spätere Implementierungen wäre auch eine direkte Kommunikation der Clients mit dem *Aggregator Daemon* denkbar. Das Script, welches auf dem Nova Controller aufgeführt wird, kann zum einen mit entsprechen Parametern als *Aggregator Daemon* gestartet werden, der die Messdaten sammelt, und zum anderen als *Balancing Daemon*, der die Daten aus der Datenbank abrufen, um die Maschinen gegebenenfalls umzuziehen. Weiterhin können dem Script Parameter übergeben werden, um die letzten Messwerte anzuzeigen sowie eine einfache Kontrolle der korrekten Funktionsweise des Systems zu ermöglichen. Zusätzlich wurde eine Funktion im Script implementiert, um die Datenbank mit Metainformationen über alle aktuell laufenden VMs zu füllen und so neue VMs automatisiert zu erfassen. Das Script überprüft darüber hinaus beim Starten eventuell fehlende Abhängigkeiten und legt gegebenenfalls eine neue Datenbank an. Das Hauptaugenmerk bei der Entwicklung bestand auf der Verwendung neuer Metriken und Energieeffizienzparameter (wie z.B. der Temperaturentwicklung) sowie deren Kombination mit bestehenden Power Management Lösungen, um eine energieeffiziente Verschiebung von VMs in aktuellen OpenStack-Installationen zu unterstützen.

## 4 Fazit und Ausblick

OpenStack hat bereits seinen festen Standplatz bzgl. des Einsatzes für Private Cloud-Umgebungen in großen IT-Infrastrukturen, wie z.B. der des CERN [Bel14], gefunden. 200 Unternehmen sind gemäß [Bel14] allein an der OpenStack-Entwicklung direkt beteiligt. OpenStack ist modular und lässt sich somit um eigene Ideen erweitern, die auch zurück in das Projekt fließen können. Ein aktueller Gegenstand der Forschung in diesem Bereich ist die adaptive Verteilung von Maschinen in der Cloud. Erste Ansätze hierfür zeigt z.B. *OpenStack Neat* [Bel13]. Während OpenStack Neat primär auf die effiziente Auslastung der Compute Nodes eingeht, wird in diesem Paper ein Prototyp vorgestellt, der OpenStack zusätzlich um Funktionen für eine adaptive Verteilung von virtuellen Maschinen

für eine Optimierung der Energieeffizienz per Live-Migration erweitert. Durch die direkte Verwendung der Standard-API und -Komponenten von OpenStack kann der Prototyp im Gegensatz zu den anderen Lösungen direkt in bestehenden OpenStack-Installationen verwendet werden. Er wird dies bezüglich in einem Testbed an der Hochschule Fulda verwendet. Hierbei sammelt der Prototyp Leistungsmetriken und kontrolliert Power Management Funktionen, die die Nodes innerhalb einer OpenStack-Umgebung zur Verfügung stellen, um eine energieeffiziente Verteilung der virtuellen Maschinen und eine Reduzierung der erforderlichen Energiekosten für Compute, Storage und Netzwerk-Komponenten in verteilten Private Clouds zu evaluieren. Neben den Metriken der Systeme kann auch die Umgebungstemperatur (z.B. in Serverschränken) gemessen werden, um diese in die optimale Verteilung von VMs einzubeziehen. Durch die Verwendung von intelligenten Steckdosen kann der Prototyp die Leistungsaufnahme und Energiekosten exemplarisch im Testbed bestimmen. Für die zukünftige Entwicklung des Prototyps wird in diesem Zusammenhang auch die Verlagerung virtueller Maschinen an andere geografische Orte, an denen z.B. die Stromkosten momentan geringer sind oder bessere klimatische Bedingungen vorliegen, bewertet. Hierbei bildet die Optimierung der Energiekosten für die erforderlichen Netzinfrastrukturen einen wichtigen Schwerpunkt. Durch die Reduzierung der erforderlichen Energie für Links und Netzwerkkomponenten zwischen kommunizierenden VMs sowie die Drosselung von aktuell nicht benötigten Redundanzsystemen, kann die Energieeffizienz in OpenStack-Umgebungen zusätzlich gesteigert werden. Zukünftig ist geplant die durch den Prototypen realisierten Erweiterungen als Plug-Ins für die für die Orchestrierung- (Heat) und Monitoring- bzw. Telemetrie-Komponenten (Ceilometer) von OpenStack umzusetzen. Durch die Verwendung des Open Cloud Computing Interfaces (OCCI) wird hierbei zukünftig zusätzlich auch eine Bewertung und Optimierung des Energiebedarfs von virtuellen Maschinen in föderierten Private und Hybrid Cloud-Lösungen über Standort- und Providergrenzen hinweg möglich.

## Literatur

- [AG13] M. A. Adnan und R. Gupta. Path Consolidation for Dynamic Right-Sizing of Data Center Networks. In *Sixth International Conference on Cloud Computing (CLOUD)*, Seiten 581–588. IEEE, 2013.
- [BB10] Anton Beloglazov und Rajkumar Buyya. Energy efficient resource management in virtualized cloud data centers. In *Proceedings of the 2010 10th IEEE/ACM International Conference on Cluster, Cloud and Grid Computing*, Seiten 826–831. IEEE Computer Society, 2010.
- [Bel13] Anton Beloglazov. Energy-Efficient Management of Virtual Machines in Data Centers for Cloud Computing. *Dissertation*, Feb. 2013.
- [Bel14] Tim Bell. Tim Bell on the Importance of OpenStack for CERN. <http://www.isgtw.org/feature/tim-bell-importance-openstack-cern/>, 2014. 11.4.2014.
- [Clo14] Rackspace Private Cloud. OpenStack Private Cloud Software Free Download Page by Rackspace. [http://www.rackspace.com/cloud/private/openstack\\_software/](http://www.rackspace.com/cloud/private/openstack_software/), 2014. 11.4.2014.

- [CRN<sup>+</sup>10] K. Christensen, P. Reviriego, B. Nordman, M. Bennett, M. Mostowfi und Juan A. Maestro. IEEE 802.3az: the road to energy efficient ethernet. *Communications Magazine, IEEE*, 48(11):50–56, 2010.
- [FLL<sup>+</sup>13] W. Fang, X. Liang, S. Li, L. Chiaraviglio und N. Xiong. VMPlanner: Optimizing virtual machine placement and traffic flow routing to reduce network power costs in cloud data centers. *Computer Networks*, 57(1):179–196, 2013.
- [Fou14a] OpenStack Foundation. OpenStack Configuration Reference. Post-Installation Configuration. <http://docs.openstack.org/havana/config-reference/content/configuring-openstack-compute-basics.html#true-live-migration-kvm-libvirt>, 2014. 11.4.2014.
- [Fou14b] OpenStack Foundation. OpenStack Manual. OpenStack Configuration Reference. Scheduling. [http://docs.openstack.org/trunk/config-reference/content/section\\_compute-scheduler.html](http://docs.openstack.org/trunk/config-reference/content/section_compute-scheduler.html), 2014. 11.4.2014.
- [GLU14] GLUSTER Incorporation. Gluster Storage Connector for OpenStack Resource Page. <http://gluster.org/community/documentation/index.php/OSConnect>, 2014. 11.4.2014.
- [HF12] Ralph Hintemann und Klaus Fichter. Energieverbrauch und Energiekosten von Servern und Rechenzentren in Deutschland. [https://www.bitkom.org/files/documents/Kurzstudie\\_\\_Borderstep\\_l\\_Rechenzentren\(1\).pdf](https://www.bitkom.org/files/documents/Kurzstudie__Borderstep_l_Rechenzentren(1).pdf), 2012. 11.4.2014.
- [HS10] Ralph Hintemann und Holger Skurk. Energieeffizienz im Rechenzentrum. In *Green-IT, Virtualisierung und Thin Clients*, Seiten 19–56. Springer, 2010.
- [HSM<sup>+</sup>10] Brandon Heller, Srinivasan Seetharaman, Priya Mahadevan, Yiannis Yiakoumis, Puneet Sharma, Sujata Banerjee und Nick McKeown. ElasticTree: Saving Energy in Data Center Networks. In *NSDI*, Jgg. 3, Seiten 19–21, 2010.
- [MG11] Peter Mell und Timothy Grance. The NIST definition of cloud computing (draft). *NIST special publication*, 800(145):7, 2011.
- [MKDK11] V. Mann, A. Kumar, P. Dutta und S. Kalyanaraman. VMFlow: leveraging VM mobility to reduce network power costs in data centers. In *NETWORKING 2011*, Seiten 198–211. Springer, 2011.
- [Ope14] OpenStack. OpenStack Open Source Cloud Computing Software. <http://www.openstack.org/>, 2014. 11.4.2014.
- [SFW<sup>+</sup>13] Aibo Song, Wei Fan, Wei Wang, Junzhou Luo und Yuchang Mo. Multi-objective virtual machine selection for migrating in virtualized data centers. In *Pervasive Computing and the Networked World*, Seiten 426–438. Springer, 2013.
- [SH13] Nongmaithem Ajith Singh und M Hemalatha. Reduce Energy Consumption through Virtual Machine Placement in Cloud Data Centre. In *Mining Intelligence and Knowledge Exploration*, Seiten 466–474. Springer, 2013.
- [VT13] Daniel Versick und Djamshid Tavangarian. CAESARA-Combined Architecture for Energy Saving by Auto-Adaptive Resource Allocation. *GI-Edition*, Seite 31, 2013.
- [WDD12] Marc Wilkens, Gregor Drenkelfort und Lars Dittmar. Bewertung von Kennzahlen und Kennzahlensystemen zur Beschreibung der Energieeffizienz von Rechenzentren. [http://opus.kobv.de/tuberlin/volltexte/2012/3573/pdf/IZE\\_3\\_wilkens.pdf](http://opus.kobv.de/tuberlin/volltexte/2012/3573/pdf/IZE_3_wilkens.pdf), 2012. 11.4.2014.

# **Virtualisierung und Cloud Services**



# FOSP: Towards a Federated Object Sharing Protocol that Unifies Operations on Social Content

Felix Maurer and Sebastian Labitzke

Karlsruhe Institute of Technology (KIT)  
Steinbuch Centre for Computing (SCC) & Institute of Telematics  
Zirkel 2, 76131 Karlsruhe, Germany  
felix.maurer@student.kit.edu, sebastian.labitzke@kit.edu

**Abstract:** Years ago, the World Wide Web (WWW) began as a system for publishing interlinked hypertext documents. While the protocols on top of which the WWW is built are almost still the same, the usage, as well as the content have changed significantly. Simple delivery of hypertext documents has been expanded by operations, such as uploading, sharing, and commenting on pieces of content. Online Social Networks (OSNs) and other IT services provide aggregated views on these pieces of content. However, the services are often implemented as vendor specific applications on top of common web technologies, such as HTTP, HTML, JavaScript and CSS. Moreover, users are locked into these applications of dedicated providers, which prevents sharing of content across applications and limits the control users have over their data. Most existing approaches that overcome these issues focus on defining a common HTTP API or prefer solutions based on peer-to-peer networks. In this paper, we start by discussing related work and identifying essential requirements for an appropriate solution. Furthermore, we outline the concept and implementation of a Federated Object Sharing Protocol (FOSP), i.e., a different approach to support today's common operations on social content already on a protocol level. We show that services built on top of this protocol can be federated by default, i.e., users registered with different providers can easily interact with each other. Finally, we provide an evaluation and discussion on the proposed approach.

## 1 Introduction

The World Wide Web started as a system to publish simple, linked documents and the Hypertext Transfer Protocol (HTTP) was created to transfer these documents. In contrast, today's websites are rather complex front ends of powerful applications and documents are dynamically generated depending on parameters of the HTTP request. Furthermore, extensions for sending information to a server were implemented and the hyperlinks initially used to connect documents became methods to invoke actions on servers.

This way, applications were created that allow a user to submit and share pieces of content with others, such as texts, pictures and videos. Nowadays, many types of websites are implemented that provide access to user generated content, e.g., weblogs, microblogging sites, and Online Social Networks (OSNs). However, many features necessary for those

platforms are not supported by HTTP and, therefore, are implemented on top of this protocol. In turn, while the protocols and data formats used are open standards, the applications are not.

As a result, users are unable to share content across platforms and can often access shared data only with a provider-specific client, although the features of the platforms are often very similar. This includes, but is not limited to, authentication and authorization, as well as adding, deleting, rating, and modifying content, or sending notifications. Furthermore, the structure of shared data shows similarities, i.e., content is often hierarchically stored, e.g., comments added to a text, video or photo. These features have to be reimplemented by each provider and the resulting implementations are largely incompatible.

In this paper, we provide the outline of a Federated Object Sharing Protocol (FOSP), i.e., a common protocol for applications that allow users to share pieces of content and operate on this content in the aforementioned manner. On the basis of functionality existing OSNs provide, this protocol allows users to register an account and log-in to a platform, share arbitrary kind of content, grant and deny access on content to certain users or groups of users and to subscribe to streams of content of others. Additionally, to prevent user lock-in with respect to a single provider, it allows federation of all platforms that implement the protocol. The novelty of this approach constitutes the handling of common operations, as well as the federation of platforms already on a protocol level. For reasons of acceptance, the protocol aims to have low complexity regarding implementation and deployment, as well as provide users with easy-to-use and familiar features.

In Section 2, we discuss existing protocols and related work. Section 3 presents common use case scenarios of OSNs and similar applications that are used to compile a list of requirements. Section 4 describes the new protocol-based approach that is evaluated and discussed in Section 5. Section 6 concludes the paper and gives an outlook on future work.

## 2 Related Work

First, we review several protocols that are widely used on the web today. Second, we examine selected open source projects that aim at supporting users in cross-platform interaction. Third, we explore scientific work that focuses on federated sharing of information.

As mentioned, the **Hypertext Transfer Protocol (HTTP)** forms the protocol basis of today's WWW. Due to its extensibility, HTTP allows the implementation of new request methods, response statuses, headers, and content types. Therefore, this protocol can be used for posting and retrieving almost all kinds of data, as well as for basic authentication. However, the semantics of provided operations is only partially defined. This results in systems that misuse operations, e.g., some applications misuse GET operations to delete a resource even though a DELETE operation exists. Furthermore, as HTTP is state-less, each request usually opens a new connection, i.e., authentication is required for each request and authorization is not handled by default. Additionally, a server cannot notify

clients about updates and while workarounds like long-polling, e.g., BOSH<sup>1</sup>, are available, it led to the development of the WebSocket protocol [FM11].

**Web-based Distributed Authoring and Versioning (WebDAV)** is an extension to HTTP that introduces new methods and semantics for specific resources [Dab09]. It implements missing functionality regarding resource manipulation for supporting authoring of web-sites. By adding meta-data to the resources, access control lists, information about resource ownerships, and further meta-information of the resources can be stored. Additional methods for manipulating meta-data allow, for instance, to list children of hierarchical structured resources. However, WebDAV is also not designed for pushing updates and to be implemented in a federated network.

**SPDY**<sup>2</sup> constitutes an extension to the way HTTP requests are transferred over TCP and serves as the working base for **HTTP 2.0** [BTP<sup>+</sup>13]. It allows multiplexing more than one request over a single connection and the server can even push resources over an existing connection. While it improves the transfer rate, it does not solve the problems of HTTP or WebDAV, like pushing notifications or federating access control.

The **eXtensible Messaging and Presence Protocol (XMPP)** [SA11] was developed as an instant messaging (IM) protocol. XMPP allows the exchange of structured messages between multiple network entities and enables the federation of servers. It is designed to be extensible, i.e., additional functionality can be defined by so called XMPP Extension Protocols (XEP). Methods for sending files, publishing the location of the user, initiating video and voice calls, and many more have been implemented. However, it is not designed for storing data and extensions are possible but lead to defining new protocols.

The **Network File System (NFS)** [SCR<sup>+</sup>03] is a protocol commonly used for distributed file systems in Unix environments. It allows accessing files over the network but does not support sending notifications about changes. The Glamor research project by IBM aims to provide a federated file system layer on top of NFS [LNT08]. It works by providing a virtual file system based on multiple NFS servers but relies on a central configuration. Therefore, it is not truly a globally federated file system.

**Wave** [WW11] is a protocol and former Google product for real time collaboration. Users work on XML-like documents called *wavelets* and changes on these documents are synchronized over the network. The protocol allows federation of servers so that users can collaborate on wavelets that are hosted by different providers. However, access control only works at the level of a single wavelet and content is primarily text based.

In addition to existing protocols, the open source community introduced many projects dedicated to free users from platform lock-in and to increase privacy<sup>3</sup>. Next to projects like StatusNet<sup>4</sup> for microblogging and SecureShare<sup>5</sup> for peer-to-peer (P2P)-based encrypted social networking, approaches to implement traditional OSNs over HTTP are invented:

---

<sup>1</sup><http://www.xmpp.org/extensions/xep-0124.html>

<sup>2</sup><http://www.chromium.org/spdy/spdy-whitepaper>

<sup>3</sup><https://gitorious.org/social/pages/ProjectComparison>

<sup>4</sup><http://status.net>

<sup>5</sup><http://secushare.org>



**Diaspora**<sup>6</sup> defines an HTTP API for federating OSN servers. The content is formatted according to the ActivityStreams specification<sup>7</sup> and transported using various HTTP-based protocols<sup>8</sup>. However, whereas Diaspora\* includes sharing of posts, comments and other content, the API is limited to certain types of content<sup>9</sup>. Additionally, instead of retrieving posts that are of interest, i.e., a user-specific or even user-defined preselection of content, the protocol simply broadcasts new posts to users who are connected with the author.

**Buddycloud**<sup>10</sup> defines an XMPP API and additionally an HTTP API. In the Buddycloud model, posts, pictures, and other files can be submitted to so called “channels” that users can subscribe to. This way, content can be shared with users even if they are registered with other providers. However, the channels restrict the way content can be organized and also restrict the types of content that can be shared. Additionally, access control only happens on the channel level, i.e., access to single posts cannot be restricted. While Buddycloud provides multiple APIs and services for OSNs it does not focus on creating a simple and versatile way of sharing data, setting access rights and sending notifications.

Besides existing protocols and open source projects, a lot of scientific work on decentralized sharing of social content has been published. Approaches range from connecting existing OSNs (cf. [KTS11, OP12]) to building new OSNs, for instance, based on P2P networks (cf. [BH13, BSVD09, BFG<sup>+</sup>10, NPA10]).

Tramp et al. introduced the **Distributed Semantic Social Network** (DSSN) [TFE<sup>+</sup>14] that exchanges data represented according the Resource Description Framework (RDF)<sup>11</sup> and uses the vocabulary of the Friend of a Friend (FOAF)<sup>12</sup>. Users can discover each other using WebID<sup>13</sup> and publish new content via the PubSubHubbub protocol<sup>14</sup>. Notifications about updates are distributed with the Semantic Pingback protocol<sup>15</sup> but are not delivered directly to the user. Access control is implemented by using WebID and by adding support for access delegation. In contrast to our approach of an integrated single protocol, the DSSN combines multiple existing protocols and data formats.

The **Distributed Platform for Multimedia Communities** of Graffi et al. is built on top of a P2P network that provides a Distributed Hash Table (DHT) [GPM<sup>+</sup>08]. The DHT is used to store distributed linked lists that contain and structure the shared content. Several services form the base of their framework: a storage and replication layer, a storage dispatcher for serialization and storage of application specific data, a message dispatcher for direct communication, etc. Different features of OSNs are then implemented as plug-ins on top of these services. To implement access control, they use asymmetric public keys per user and symmetric keys to encrypt content shared with a group of users. However, a general mechanism to receive notifications about new or changed content is missing.

---

<sup>6</sup><https://diasporafoundation.org>

<sup>7</sup><http://activitystrea.ms/>

<sup>8</sup><http://www.salmon-protocol.org/>, <http://tools.ietf.org/html/draft-ietf-appsawg-webfinger-18>

<sup>9</sup>[http://wiki.diasporafoundation.org/Main\\_Page](http://wiki.diasporafoundation.org/Main_Page)

<sup>10</sup><http://buddycloud.com>

<sup>11</sup><http://www.w3.org/RDF/>.

<sup>12</sup><http://www.foaf-project.org/>

<sup>13</sup><http://www.w3.org/wiki/WebID>

<sup>14</sup><http://code.google.com/p/pubsubhubbub/>

<sup>15</sup><http://aksw.org/Projects/SemanticPingback.html>

| Project                                         | Architecture | # Components            | Protocols                                     |
|-------------------------------------------------|--------------|-------------------------|-----------------------------------------------|
| Diaspora*                                       | Federated    | 1 server                | HTTP (Salmon, Webfinger)                      |
| Buddycloud                                      | Federated    | 3 services              | HTTP, XMPP                                    |
| DSSN                                            | Federated    | 1 and more services     | HTTP (WebID, PubSubHubbub, Semantic Pingback) |
| Distributed Platform for Multimedia Communities | P2P          | 1 (+ framework plugins) | DHT (PAST <sup>17</sup> )                     |
| SODESSON                                        | P2P          | 4 framework components  | DHT                                           |
| Safebook                                        | P2P          | 1 client + 1 service    | DHT                                           |

Figure 1: Comparison of OSN projects

Another P2P-based approach constitutes **SODESSON** [BH13] introduced by Baumgart et al. SODESSON abstracts from users' devices to be able to directly address the users and is focused on providing services directly on (mobile) devices while using the social graph as basis for access control decisions. The framework is divided into multiple components, such as the connection manager, the distributed data storage, the contact manager, and the service manager that constitutes the interface to the actual applications. Services that can be provided are divided into three types, namely direct services between online devices, persistent services stored in a distributed storage and hybrid services as a mix of both. Similar to our goals, SODESSON enables distributed sharing of data in OSNs but specializes on providing a transparent access to services on multiple user devices.

Cuttillo et al. introduced **Safebook**, i.e., a P2P-based OSN, too [CMS09]. It is focused on user privacy and the utilization of multiple techniques to prevent information leakage on the application layer and on the transport layer as well. In Safebook, each user has a layered network of peers that forward requests, similar to the Tor network<sup>16</sup>. The most inner layer of peers that is directly connected to the user also mirrors profile information and comments, so that they are available even if the user is offline. To prevent eavesdropping, all communication is encrypted. While it provides methods for publishing and retrieving profile information, it lacks the possibility of subscribing and sending notifications and is limited to certain types of content. Furthermore, an additional service is needed for creating a new identity and uses out of band mechanisms to verify it.

In summary, existing protocols satisfy a part of the requirements for federated sharing of social content (cf. Section 3). However, none of them constitutes a "Jack of all trades" that, additionally, can be easily deployed. HTTP/WebDAV and NFS work well for storing data and handling access control, whereas XMPP and Wave support federation and include notifications. Community projects focus on a smaller set of features and prioritize working software. Diaspora\* uses multiple data formats and does not include notifications while Buddycloud uses multiple protocols. Most scientific work prefer solutions on top of P2P networks and provide security and privacy mechanisms (cf. Fig. 1).

<sup>16</sup><https://www.torproject.org/>

### 3 Use Cases and Requirements

In this section, we start by introducing several use cases to evaluate a new protocol with respect to its applicability in real world scenarios. The use cases are derived from functionality provided by existing OSNs extended by federated sharing scenarios.

**Use cases:** Independently of the provider a user is registered with, she can share content with users of other providers, as long as they support the API or protocol (**UC1**). Users can subscribe to new content of other users and will receive notifications about added or changed content (**UC2**). Users can comment on or “like” content, if allowed by the content author (**UC3**). Users can publish personal information about themselves, including but not limited to age, sex, location and personal preferences (**UC4**). Users should not be limited in what kind of content they can share (**UC5**). Users can restrict access to their shared content on a per-user-basis or on user-defined groups (**UC6**).

Based on these use cases, we derive requirements that must be met by the approach introduced in this paper to allow a simple solution while still providing the necessary features:

**Federation (R1):** To prevent platform lock-in, the federation of servers must be possible. Each server must be able to become part of the global network just by implementing the protocol. Therefore, all users must have a globally unique identifier and their content must also have globally unique addresses.

**Access control model (R2):** As users want to limit access to their content, they must be able to define access rights. Therefore, a discretionary access control (DAC) [SV01] policy should be preferred. Adding role based access control (RBAC) [SV01] is desirable, as users likely sort their peers into groups and then grant or deny access to whole groups.

**Authentication and authorization (R3):** To enable access control, servers must be able to authenticate and authorize users. Furthermore, as each server can only authenticate users that are registered with itself, it has to act as an identity provider for other servers.

**Content manipulation (R4):** Users must have multiple methods available to interact with content. This includes adding new, as well as changing and deleting existing content. The type of content that can be shared must not be restricted. Furthermore, the user must be able to identify content she wants to manipulate and, therefore, each content unit must be addressable by a URI.

**Subscription (R5):** A user must be able to express interest in (a set of) specific content and will be notified if content is added or altered. The information about these subscriptions must be stored as meta-data, such as access control lists.

### 4 The Federated Object Sharing Protocol (FOSP)

Subsequently to the presentation of related work and requirements regarding federated sharing of social content, we provide another contribution in the following, i.e., we outline the concept and implementation of a new protocol for federated data sharing.

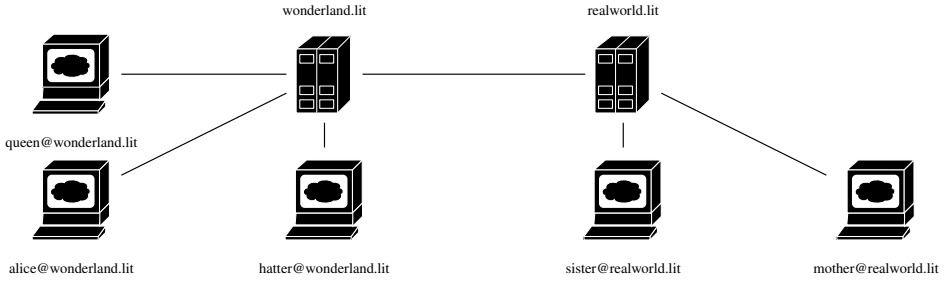


Figure 2: The network architecture of FOSP. User Alice connects to her *home server* wonderland.lit

**Network Architecture:** In a FOSP network, each agent is either a client or a server. Each server belongs to one provider identified by a domain name. Users connect via a client to the server of their provider that we call the user’s *home server*. To retrieve or alter content, the client sends requests. If the *home server* is not responsible for the requested content, it connects to the server that actually stores the content and relays the request (cf. Fig. 2).

**Data Structure:** We refer to the basic unit of data as an *object* that consists of key-value pairs with specific meanings (cf. “Policies”) and adheres to the JSON [Cro06] specification. Furthermore, a file can be attached to an object and all objects are stored as nodes within trees. Each user owns one tree of objects stored on her home server and the root of the tree is named by her globally unique identifier.

**Policies:** A set of policies define how the content of the objects is to be interpreted and the servers must enforce these policies. Some of the key-value pairs store meta-data, such as the owner of an object, the time it was created, or information about an attached file. An “acl” field of an object stores the associated access control list. A server then determines whether or not a user can perform certain operations on the object. A “subscription” field stores information about users that want to be notified on changes of an object or one of its descendants, which is evaluated by the server responsible for this object. An example for the default key-value pairs of an object is shown in Fig. 3. Furthermore, whether a server can relay a message or accept a relayed message is also subject to policies. For example, a server must only accept a relayed request if the domain name of the remote server matches the domain name in the identifier of the user the request originates from.

**Messages:** The messages that are sent between the servers and clients are divided into three different types. Clients send *requests* to retrieve or manipulate objects and requests have request types, i.e., SELECT, UPDATE, or CREATE that determine how the request acts upon an object. As an answer to a request, servers send *responses* that also have a type, which is either SUCCEEDED or FAILED. Furthermore, if a request changes objects, the server sends *notifications* to users that have subscribed to these changes and the notification includes an event type, i.e., UPDATED that corresponds to the type of change. All messages carry different information in their first line, depending on their type. Additionally, like in HTTP, each message can have headers that modify its meaning and a body that carries the payload. Messages are serialized as UTF-8 text, except for the body, which consists either also of UTF-8 text or of arbitrary binary data when sending files.

```
{ btime: "2007-03-01T13:00:00Z",
  mtime: "2008-05-11T15:30:00Z",
  owner: "alice@wonderland.lit",
  acl: {
    owner: ["read-data", "write-data", "read-acl", "write-acl",
            "read-subscriptions", "write-subscriptions", "read-children",
            "write-children", "delete-children"],
    users: { alice@wonderland.lit: [ "read-data", "write-data", ... ] }
  },
  subscriptions: {
    users: { alice@wonderland.lit: { events: [ "created", "updated" ], depth: 1 } }
  },
  type: "text/plain",
  data: "Just plain text" }
```

Figure 3: An example object owned by alice@wonderland.lit.

The format of a message looks similar to an HTTP request. Each message is then send as an WebSocket message over a WebSocket connection.

**Implementation and performance evaluation:** Based on the presented concept, we built a proof-of-concept implementation of a FOSP server and two FOSP clients both written in JavaScript [Mau14]. The server is written for the node.js<sup>18</sup> runtime environment and uses the RethinkDB<sup>19</sup> as database. One of the clients is implemented as a single page JavaScript application<sup>20</sup> and highlights the ease of deployment that can be achieved. The other client is a command line client that can be used in a shell. For performance evaluation, we ran several tests with two instances of the server and up to 800 user clients. While the server already achieves fast response times and can handle many consecutive requests per minute, it is not yet optimized for many parallel users due to the chosen programming language and database. However, FOSP as a protocol can be implemented with almost any language and backing database. We already work on a Go<sup>21</sup> and PostgreSQL<sup>22</sup> implementation that performs significantly better on concurrent requests, which we plan to make public for download. Furthermore, while scaling the server of one provider might be more difficult, scaling by adding providers should be trivial, similar to XMPP and Buddycloud.

## 5 Evaluation and Discussion

In Section 3, we defined several use cases and requirements that should be supported by a new protocol for federated sharing of social content. In the following, we show that FOSP meets all of these criteria. The federation of different networks (UC1 and R1) is achieved by relaying messages between servers that, additionally, act as identity providers for their users (R3). Users can share arbitrary content by encoding it as JSON and storing it in an object or by saving it as a file attached to an object (UC5 and R4). They can also subscribe to changes on an object and its descendants and will receive notifications about added or

<sup>18</sup><http://nodejs.org/>

<sup>19</sup><http://www.rethinkdb.com/>

<sup>20</sup><http://singlepageappbook.com/>

<sup>21</sup><http://golang.org/>

<sup>22</sup><http://www.postgresql.org/>

altered content (UC2 and R5). Comments can be added as new child objects to existing objects (UC3). Personal information can also be saved in an object or multiple objects to allow more granular access control (UC4). Furthermore, an access control list on each object allows users to share content with a group of peers or even a single user. Access control is enforced by the server that stores the object and users can grant or deny rights by setting corresponding entries in the access control list of the object (UC6, R2, and R3).

In this paper, we introduced the first steps towards a new protocol for federated sharing of social content. However, some limitations remain that require further work on this project to reach a protocol that can be globally used. First, server to server authentication must work reliably in order to prevent data leakage in the federated network. This can be done by using DNS but also with more sophisticated methods, such as dial-back strategies or SSL certificate verifications. Furthermore, if a user grants access to content to another user of a different provider, the access is implicitly granted to this provider as well.

However, despite open questions, FOSP could already be used today in specific scenarios in which trust across participating servers is implicitly given but central infrastructures are not deployable. For instance, sharing content between members of universities each of which operates one trustworthy FOSP server would be possible by use of the FOSP version presented within this paper (cf. necessary trust between identity providers of authentication and authorization infrastructures, such as within the DFN-AAI<sup>23</sup>). Furthermore, FOSP constitutes a basis for further discussion and research on a path completely different compared to existing approaches, i.e., a new protocol for provider independent content sharing. Moreover, existing approaches can be combined with FOSP and variations or extensions are imaginable, e.g., just with slight adjustments, FOSP could be used in P2P networks or objects could be extended to allow encrypted content, versioning or locking.

## 6 Conclusion and Future Work

In this paper, we analyzed related work and identified use cases and requirements for federated sharing of social content in order to propose an approach for implementing today's common operations already on a protocol level. Furthermore, we introduced the concept and implementation of the Federated Object Sharing Protocol (FOSP) that is designed to be simple and deployable while meeting the identified requirements and use cases. FOSP combines storing data, access control and publish-subscribe mechanisms for future content sharing. Furthermore, it enables federation of servers that belong to different providers by default. This way, FOSP can simplify the implementation of traditional social networks and allows competing but compatible server and clients. In the future, we will refine the FOSP specification into a document that covers all implementation relevant details and will add further features, such as content encryption that can remedy remaining security problems. We also plan to publicly provide our ready-to-use FOSP client and FOSP server implementation that can handle a large amount of users for download.

---

<sup>23</sup><https://www.aai.dfn.de/>.

## References

- [BFG<sup>+</sup>10] Marin Bertier, Davide Frey, Rachid Guerraoui, Anne-Marie Kermarrec, and Vincent Leroy. The Gossple Anonymous Social Network. In *Middleware 2010*, volume 6452. Springer Berlin Heidelberg, 2010.
- [BH13] I. Baumgart and F. Hartmann. User-centric networking powered by SODESSON. *PIK - Praxis der Informationsverarbeitung und Kommunikation*, 36(2), May 2013.
- [BSVD09] Sonja Buchegger, Doris Schiöberg, Le-Hung Vu, and Anwitaman Datta. PeerSoN: P2P social networking: early experiences and insights. In *Proc. of the 2nd ACM EuroSys Wksp. on Social Network Systems (SNS'09)*, New York, NY, USA, 2009. ACM.
- [BTP<sup>+</sup>13] M. Belshe, Twist, R. Peon, Inc Google, M. Thomson, Microsoft, A. Melinkov, and Isode Ltd. Hypertext Transfer Protocol version 2.0, August 2013.
- [CMS09] L.A. Cutillo, R. Molva, and T. Strufe. Safebook: A privacy-preserving online social network leveraging on real-life trust. *Communications Magazine, IEEE*, 47(12), 2009.
- [Cro06] D. Crockford. The application/json Media Type for JavaScript Object Notation (JSON). RFC 4627 (Informational), July 2006.
- [Dab09] C. Daboo. Extended MKCOL for Web Distributed Authoring and Versioning (Web-DAV). RFC 5689 (Proposed Standard), September 2009.
- [FM11] I. Fette and A. Melnikov. The WebSocket Protocol. RFC 6455 (Proposed Standard), December 2011.
- [GPM<sup>+</sup>08] K. Graffi, S. Podrajanski, P. Mukherjee, A. Kovacevic, and R. Steinmetz. A Distributed Platform for Multimedia Communities. In *Proc. of the 10th IEEE Int'l Symp. on Multimedia (ISM'08)*. IEEE, 2008.
- [KTS11] Moonam Ko, H. Touati, and M. Shehab. Enabling Cross-Site Content Sharing between Social Networks. In *Proc. of the 3rd Int'l Conf. on Privacy, Security, Risk and Trust (PASSAT'11) and on Social Computing (SOCIALCOM'11)*. IEEE, 2011.
- [LNT08] U. Lanjewar, M. Naik, and R. Tewari. Glamor: An architecture for file system federation. *IBM Journal of Research and Development*, 52(4.5), 2008.
- [Mau14] Felix Maurer. FOSP: Federated Object Sharing Protocol. Bachelor's thesis, Karlsruhe Institute of Technology (KIT), Dept. of CS, Inst. of Telematics, January 2014.
- [NPA10] R. Narendula, T.G. Papaioannou, and K. Aberer. Privacy-Aware and Highly-Available OSN Profiles. In *Proc. of the 19th IEEE Int'l Wksp. on Enabling Technologies: Infrastructures for Collaborative Enterprises (WETICE'10)*. IEEE, 2010.
- [OP12] Alexandra Olteanu and Guillaume Pierre. Towards robust and scalable peer-to-peer social networks. In *Proc. of the 5th Wksp. on Social Network Systems (SNS'12)*, New York, NY, USA, 2012. ACM.
- [SA11] P. Saint-Andre. Extensible Messaging and Presence Protocol (XMPP): Core. RFC 6120 (Proposed Standard), March 2011.
- [SCR<sup>+</sup>03] S. Shepler, B. Callaghan, D. Robinson, R. Thurlow, C. Beame, M. Eisler, and D. Noveck. Network File System (NFS) version 4 Protocol. RFC 3530 (Proposed Standard), April 2003.
- [SV01] Pierangela Samarati and SabrinaCapitani Vimercati. Access Control: Policies, Models, and Mechanisms. In *Foundations of Security Analysis and Design*, volume 2171. Springer Berlin Heidelberg, 2001.
- [TFE<sup>+</sup>14] Sebastian Tramp, Philipp Frischmuth, Timofey Ermilov, Saeedeh Shekarpour, and Sören Auer. An Architecture of a Distributed Semantic Social Network. *Semantic Web*, 5(1), 2014.
- [WW11] T. Weis and A. Wacker. Federating Websites with the Google Wave Protocol. *Internet Computing, IEEE*, 15(3), 2011.

# Der Weg von BYOE zum GYSE

Christopher Ritter, Patrick Bittner, Odej Kao

tubIT – IT-Service-Center der TU-Berlin  
Technische Universität Berlin  
Einsteinufer 17  
10587 Berlin

[christopher.ritter@tu-berlin.de](mailto:christopher.ritter@tu-berlin.de)

[patrick.bittner@tu-berlin.de](mailto:patrick.bittner@tu-berlin.de)

[odej.kao@tu-berlin.de](mailto:odej.kao@tu-berlin.de)

**Abstract:** Mit der zunehmenden IT-Unterstützung von Studium und Lehre unterliegen die Universitätsrechenzentren einem Paradigmenwechsel hin zum IT-Service-Center. Das Angebot an IT-Diensten und digitalen Medien im Bereich des Studiums unterliegt derzeit einem stetigen Wachstum. Durch BYOD (Bring Your Own Device) wurde dieser Wandel noch beschleunigt. Mit dem derzeitigen Umstieg von BYOD zum BYOE (Bring Your Own Environment) führt die Vielzahl an neuen IT-Diensten aber nicht nur zu einer Erleichterung für die Studierenden. Viele der Digital Natives verfügen zum Zeitpunkt der Immatrikulation nicht nur über eigene IT-Ausstattung, sondern haben auf dieser bereits alle von ihnen benötigten Dienste installiert und konfiguriert. Die Realisierung eines zentralen, umfassenden Identitätsmanagementsystems und der damit verbundenen Unterstützung von Single Logon und Single SignOn ermöglicht zwar die Nutzung der Dienste mit einer einzigen Benutzerkennung, in der existierenden Infrastruktur werden weitere, meist bereits in anderer Form vorhandene Dienste aber häufig als störend empfunden. Einige für das Studium benötigte Anwendungen sind unter Umständen nicht mit der vorhandenen Konfiguration kompatibel oder führen gar zu Störungen der gewohnten Umgebung. Die Erwartung an die IT der Universität ist, neben der Integration der eigenen, bestehenden IT-Umgebung in die von der Universität bereitgestellten Systeme für Studium und Verwaltung, unabhängig vom Provider oder der spezifischer Ausprägung, daher auch die Möglichkeit eine auf ihr Studium abgestimmte Umgebung auf ihren Geräten nutzen zu können.

Diese Aufgabe ist von den Universitäten zu leisten, erfordert jedoch ein frühzeitiges Umdenken und Flexibilität an der zentralen Stelle des Campus Management: bei den Systemen für Identity und Service Management. Diese müssen BYOE bereits bei der Provisionierung unterstützen und als Normalfall betrachten. Darüber hinaus müssen die für das Studium des Studierenden angebotenen Anwendungen und Medien gebündelt und als eine Umgebung zur Verfügung gestellt werden können.

In dem folgenden Beitrag wird ein System beschrieben, das als erster Ansatz zur Lösung der identifizierten Probleme dienen soll. Ausgehend von der nahtlosen



Erfassung aller Mitglieder der Universität und deren aktueller Kontexte, wird eine Arbeitsumgebung generiert, die sowohl bezogen auf enthaltene Dienste als auch auf die zur Verfügung stehenden Inhalte explizit auf die Bedürfnisse des Studierenden angepasst ist. Der aktuelle Stand der Entwicklung umfasst bisher eine Basisplattform mit entsprechender Grundfunktionalität, die in weiteren Iterationen sukzessive ausgebaut werden muss.

## 1 Einleitung

Während die Generationen vor den Digital Natives zum Zeitpunkt der Immatrikulation nur teilweise über eine eigene IT-Ausstattung (sowohl Hardware als auch Dienste) verfügten und folgerichtig von der Universität die Bereitstellung aller notwendigen IT-Komponenten für ein erfolgreiches Studium (E-Mail, Speicher, Webseite, WLAN, Verwaltungssysteme, ...) erwarteten, änderte sich diese Anforderung zunächst mit dem Aufkommen von BYOD (Bring Your Own Device).

Mit stark wachsender Zahl begannen die Studierenden bereits mit ihrer eigenen Hardware ihr Studium. Mittlerweile erleben die Universitäten den Umstieg vom BYOD zum BYOE (Bring Your Own Environment). Die aktuellen Generationen kommen komplett ausgestattet an die Universität. Dies betrifft nicht mehr nur bereits vorhandene Hardware, sondern auch eine vollständige Infrastruktur an Diensten inkl. vorhandener E-Mailkonten und mehreren Konten in sozialen Netzwerken.

Diese Ausstattung wird in Teilen bereits vor dem Studium vorausgesetzt, da schon der erste Schritt – die Onlinebewerbung um einen Studienplatz – eine funktionierende IT-Umgebung voraussetzt. Die zahlreichen frei verfügbaren IT-Dienste sowie der deutlich frühere Einstieg in die elektronische Kommunikation als vor Jahren üblich, führen dazu, dass die Studierenden gar keine neue – universitäre – IT-Umgebung haben wollen, die ihre alte Umgebung ersetzt. Sie sehen ihre eigene IT-Umgebung als lebenslang begleitend, die auch nach den 5-6 Jahren Universitätsaufenthalt bestehen bleiben wird. Häufig werden das zunehmende Angebot an IT-Diensten und die bestehende Notwendigkeit diese zu nutzen als störend empfunden. Zum einen liegt dies in der Tatsache begründet, dass man sich bereits an die Nutzung von Diensten mit ähnlichem oder gar gleichem Funktionsumfang gewöhnt hat, zum anderen an dem Umstand, dass einige Dienste in ihrer Bedienungsführung dem gewohnten Nutzerverhalten entgegen stehen.

Besondere Ablehnung gegen die angebotenen Dienste entsteht in den Fällen, in denen im Rahmen des Studiums spezielle Software genutzt werden muss und diese gegebenenfalls Inkompatibilitäten mit dem eigenen System aufweist, die unter Umständen sogar Auswirkungen auf das restliche System haben.

Neben den angebotenen Diensten wächst aber auch die Anzahl an Medien, wie Vorlesungsskripten, Büchern, Vorträgen etc. die in digitaler Form bereitgestellt werden. Diese sind häufig an unterschiedlichen Stellen zu finden.

Die Erwartung an die IT der Universität endet daher nicht bei der Integration der eigenen, bestehenden IT-Umgebung in die offiziellen Systeme für Studium und Verwaltung, unabhängig vom Provider oder spezifischer Ausprägung, sondern geht hin zu einer Studiums bezogenen Umgebung, die bereits alle relevanten Anwendungen installiert und konfiguriert hat, Zugriff auf benötigte Medien bietet und auf dem eigenen System ausgeführt werden kann.

Um diese Aufgabe leisten zu können, ist von den Universitäten ein frühzeitiges Umdenken sowie Flexibilität an der zentralen Stelle des Campus Management, den Systemen des Identity- und Service Management, unausweichlich. Diese müssen BYOE bereits bei der Provisionierung unterstützen und als Normalfall betrachten. Die Dienste müssen – jedenfalls für die nicht datenschutzrelevanten Sachverhalte – Providerneutral angeboten werden. Inhalte müssen gebündelt und Studiums bezogen verfügbar gemacht werden. Dabei muss dem Studierenden die Auswahl der benötigten Dienste und Inhalte erleichtert werden und möglichst zentral verwaltbar sein.

Als Basis für die Umsetzung dieser Aufgaben wird an der TU Berlin das Identity Management System TUBIS eingesetzt. Das personalisierte Dienstportal verfolgt den Ansatz einer rollenbasierten Zugangsregelung. Jedes Mitglied der TU Berlin wird automatisch erfasst und mit Standardrollen ausgestattet, die im Arbeitsalltag durch Delegation, Stellvertretung oder Übertragung von Funktionen um weitere Rollen ergänzt werden können. Das System wird von mehr als 37000 Mitgliedern der TU Berlin seit mehr als 5 Jahren täglich genutzt.

Als zweite zentrale Komponente hat die TU-Berlin den Cloudspeicherdienst ownCloud aufgesetzt und stellt diesen flächendeckend mehr als 37000 Mitgliedern zur Verfügung.

Der Rest des Beitrags ist wie folgt strukturiert:

Kapitel 2 bietet einen groben Hintergrund über das an der TU-Berlin entwickelte Identitätsmanagementsystem „TUBIS“, während im Kapitel 3 ein Framework vorgestellt wird, mit dem eine Verknüpfung des IDM mit einem Cloudspeicherdienst sowie dessen Integration in den Prozess des Provisioning realisiert werden kann. Das letzte Kapitel gibt abschließend einen Ausblick über die weiterführende Entwicklung.

## **2 Identitätsmanagement als Grundlage**

Die Veränderungen der letzten Jahre in der IT-Hochschullandschaft spiegeln den Trend zur Integration vielfältiger Universitätsbereiche wieder, die vorher weitgehend unabhängig voneinander gearbeitet haben. Der Bedarf für eine solche Integration entsteht durch die Erwartungen von Studierenden und Mitarbeitern, zahlreiche Dienste durch Selbstbedienungsfunktionen zu nutzen und somit mehr Zeit für Studium und Forschung zu gewinnen. Die Umsetzung eines integrierten Dienstangebots stellt die Hochschulen wiederum vor signifikante technische und organisatorische Herausforderungen. Geschäftsprozesse und Verantwortlichkeiten sind selten umfassend

dokumentiert, wodurch eine übergreifende Planung, Steuerung und Operationalisierung erschwert wird.

Jedes Mitglied der Universität muss eindeutig und mit allen Befugnissen, Kontexten wie etwa Status (Studierender, Mitarbeiter), Rolle (Dekan, FG-Leiter, Abteilungsleiter, ...) oder Studiengang, bekannt sein, damit ein möglichst einfacher, einheitlicher Zugang zu allen für ihn relevanten Diensten möglich wird; und dies gleichgültig, ob sich die Person am Arbeitsplatz, zu Hause oder auf einer Dienstreise befindet.

Hierbei reifte in den letzten Jahren zunehmend die Erwartungshaltung, eigene, private Endgeräte zur Erfüllung der gestellten Aufgaben nutzen zu können. Dabei konnte es sich sowohl um leistungsstarke Notebooks als auch um leistungsbegrenzte Smartphones oder Tablets handeln. Dies führte zu einer verstärkten Verlagerung der Dienste hin zu webbasierten Angeboten.

Die TU-Berlin hat zur Erfüllung der Aufgaben eines Identitätsmanagementsystems das universitätsweite, rollenbasierte Autorisierungssystem TUBIS entwickelt [HKR08b, HR07]. Dieses integriert sich in die bestehende Infrastruktur der Campusmanagementsysteme und bietet den jeweiligen Einrichtungen die Möglichkeit der dezentralen Rechteverwaltung [RHK10]. Einen zentralen Baustein des Identitätsmanagementsystems bildet dabei eine Reihe von webbasierten Selbstbedienungsfunktionen. Durch die Plattformunabhängigkeit bildeten diese bereits die Voraussetzung zur besseren Unterstützung von BYOD. Durch gezielte Erweiterungen wurde das Paradigma des GYSE noch weiter unterstützt.

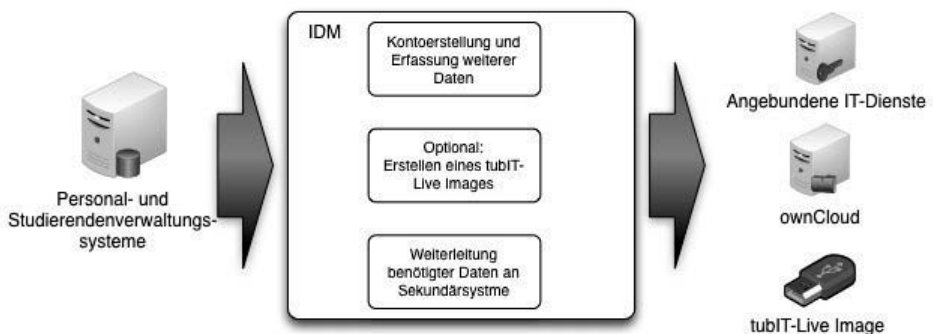


Abbildung 1 Auszug aus dem Provisioning

Zentraler Bestandteil des Identitätsmanagementsystems ist das Provisioning. Dieses ermöglicht das frühzeitige Erfassen eines neuen Mitglieds sowie dessen zentrale Versorgung mit allen benötigten Zugangsdaten und Rechten. Dieser Prozess wird auch im nachfolgend vorgestellten Framework eine zentrale Rolle spielen.

### 3 GYSE

Ziel war es, den Studierenden ein System zur Verfügung zu stellen welches ihnen ermöglicht, alle für ihr Studium relevanten Dienste zu verwenden. Das System sollte außerdem in der Lage sein, übersichtlichen Zugang zu allen für das Studium benötigten Inhalten (z.B. Vorlesungsskripte, Anleitungen, Veröffentlichungen) zu schaffen. Lange Installationsroutinen und Änderungen an der bestehenden IT-Infrastruktur der Studierenden sollten dabei vermieden werden. Das an der TU-Berlin hierfür entwickelte System teilt sich in zwei Bereiche auf: Dem Erstellen eines flexiblen OS-Images, auf welchem bereits alle benötigten Dienste installiert und konfiguriert sind, sowie die Verwaltung der für den Studenten verfügbaren Inhalte. Beide Bereiche werden dabei über das Identitätsmanagementsystem verwaltet und synchronisiert.

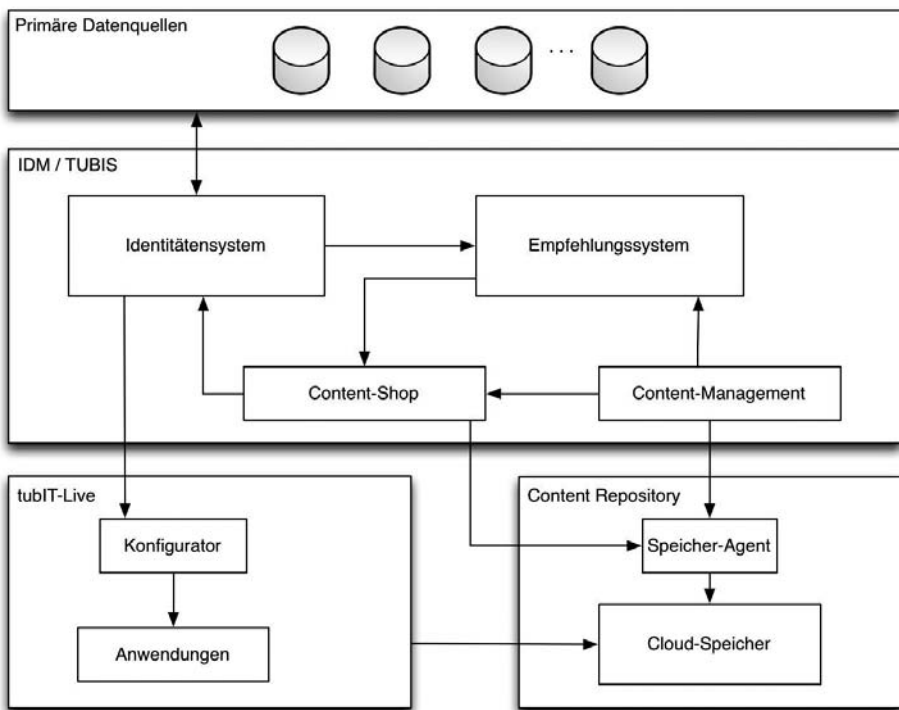


Abbildung 2 : GYSE Framework

Um sicherzustellen, dass das System allen Studierenden angeboten werden kann, wird ein zentraler Einstiegspunkt verwendet: Die Provisionierung.

Während dieses Prozesses registriert sich der Studierende mit seinen Daten, erstellt ein Benutzerkonto und wählt dafür ein persönliches Passwort aus. Zur Vermeidung redundanter oder fehlerhafter Daten, werden die Stammdaten der Studierenden für die Provisionierung direkt aus den Systemen der Studierendenverwaltung geladen.

Die meisten neu immatrikulierten Studierenden besitzen beim Eintritt in die TU-Berlin bereits ein privates E-Mailkonto. Um dem Studierenden die Möglichkeit zu geben, seine eigenen Dienste weiterzuverwenden, kann er während der Provisionierung entscheiden, ob er die Erstellung eines weiteren Postfaches wünscht oder stattdessen lediglich ein TU-Alias erstellt werden soll, um studienrelevante Nachrichten an ein bestehendes E-Mailpostfach versenden zu können. Im letzten Schritt der Provisionierung stehen dem System nun alle benötigten Daten wie z.B. der Kontoname, das Passwort und der Studiengang zur Verfügung. Diese Daten können verwendet werden, um ein personalisiertes System zu erstellen. Nachfolgend werden dem Studierenden die verfügbaren Inhalte angezeigt. Innerhalb eines „Content-Shop“ des IDM-Systems, welches an einen Onlineshop angelehnt ist, kann ausgewählt werden, zu welchen Inhalten ein Zugang gewünscht ist. Neben den zwei Kategorien Daten und Software unterteilt sich der Content-Shop in zwei große Bereiche. Während der eine Bereich alle zur Verfügung stehenden Inhalte, sortiert nach Studiengängen, auflistet, liefert der zweite Bereich eine speziell auf den Studierenden zugeschnittene Auswahl an Inhalten. Diese Vorauswahl dient dazu, die Übersicht zu verbessern und kann vom Studierenden nach eigenen Wünschen angepasst werden. Die Auswahl wird dabei von einem Empfehlungssystem generiert, welches vorhandene Daten des Studierenden verwendet, um eine repräsentative Vorauswahl der Inhalte zusammenzustellen. Vorerst werden dabei nur die besuchten Studiengänge sowie der jeweiligen Fachsemester berücksichtigt. Das System unterstützt aber bereits das Heranziehen aller im IDM verwalteten Informationen zur Generierung einer Empfehlung. Die Verwaltung modelunabhängiger Metadaten ist bereits in Planung.

### **3.1 Das Live-System**

Das OS-Image soll am Ende der Provisionierung automatisch erstellt werden. Dies wird durch einen Appliance-Buildserver realisiert, welcher über eine Web-Service Schnittstelle angestoßen wird. Aktuell wird als Basisbetriebssystem ein modifiziertes Ubuntu-Derivat verwendet. Über die mitgelieferten Systemwerkzeuge wie yum kann dieses unabhängig vom Content-Shop erweitert werden. Ein im erzeugten System integrierter Konfigurator ist für die personenbezogene Einrichtung der installierten Dienste wie E-Mailclient, Webbrowser oder den Zugang zum AFS oder zum Cloud-Speicher zuständig. Dazu verwendet der Konfigurator die übermittelten Personendaten aus der Provisionierung.

Das Image wird in einem Format zur Verfügung gestellt, welches es dem Studierenden ermöglicht, es entweder direkt von einem USB-Stick aus zu starten, oder als virtuelle Maschine innerhalb einer Virtualisierungssoftware zu betreiben. Nach der Erstellung steht das Image dem Studierenden drei Tage lang in verschlüsselter Form zum Download zur Verfügung und wird anschließend von den Servern der TU-Berlin entfernt. Wird vom Studierenden z.B. das Passwort geändert, so wird das System neu erstellt (unter Berücksichtigung des neuen Passwortes) und steht erneut für eine bestimmte Zeit zum Download bereit. Auf diese Weise sind die konfigurierten Dienste

weiterhin nutzbar. Dieser Prozess kann auch ohne Passwortänderung manuell vom Studierenden angestoßen werden. Alternativ kann das fertige Image auch direkt im Cloud-Speicher abgelegt werden und von dort auf das gewünschte Zielgerät heruntergeladen werden. Die Zugriffsrechte werden dabei durch das IDM gesteuert.

Nutzer, die manuelle Anpassungen am System vorgenommen haben und ihr Passwort ändern sind nicht gezwungen das neu erzeugte Image zu nutzen. Alternativ kann der Konfigurator im System erneut aufgerufen werden um das neue Passwort zu registrieren.

### **3.2 Zugriff zu den Inhalten**

Daten die während des Arbeitens mit dem persönlichen System entstehen oder verändert werden, müssen auch über die Neuerstellung des Systems hinaus Bestand haben. Um dies zu erreichen wird auf dem System bereits der Zugang zu einem Cloudspeicher-Dienst eingerichtet.

Jedes Mitglied der TU-Berlin erhält mit der Provisionierung auch einen Zugang im Cloud-Speicher mit einer persönlichen Quota abhängig davon, ob es sich um einen Mitarbeiter oder einen Studierenden handelt. Neben den persönlichen Daten werden hier auch die studienrelevanten Inhalte abgelegt. Um die Quota der Studierenden nicht mit den ausgewählten Inhalten zu belasten, werden keine Kopien der Daten angelegt, sondern lediglich Zugriffe auf verteilte Ordner verwaltet.

Das IDM-System steuert dabei, wer Zugriff auf welchen Bereich erhält und unterscheidet zwischen zwei Benutzergruppen.

Die Nutzer der Inhalte, also in erster Linie die Studierenden, erhalten lesenden Zugang zu all den Ordnern, die sie im Content-Shop des IDM ausgewählt haben. Diese Auswahl kann jederzeit über das persönliche Portal angepasst werden.

Die Bereitsteller der Inhalte erhalten vom IDM schreibenden Zugriff auf alle Ordner, für die sie eine entsprechende Verwalterrolle besitzen. Über eine Selbstbedienungsfunktion im persönlichen Portal können diese Nutzer auch weitere verteilte Ordner im Cloud-Speicher anlegen und mit entsprechenden Meta-Daten versehen.

Beide Benutzergruppen haben den Vorteil, dass sie von jedem System aus, für das ein Cloud-Speicher-Client zur Verfügung gestellt wird, Zugriff auf die Daten haben.

## **4 Ausblick**

Das entwickelte Framework ist ein Grundgerüst, welches in unterschiedlicher Weise weiterentwickelt werden kann. Die einzelnen vorgestellten Bereiche decken bereits die Grundfunktionalitäten ab, welche erweitert werden können, um den Benutzern so einen größeren Komfort zu bieten.

So kann die Benutzung des Content-Shops als zwingender Bestandteil des Provisioning-Prozesses etabliert werden, was dafür sorgen würde, dass jeder Benutzer (da jeder Benutzer provisioniert sein muss) ein entsprechendes Live-System besitzt.

Ausgehend vom beschriebenen Grundgerüst, kann die Integration von Software in das OS-Image weiter untersucht werden. Diese Software kann direkt bei Erstellung des Images integriert werden. Dabei ist darauf zu achten, dass die (im Content-Shop ausgewählte) Software mit anderen ausgewählten Softwarepaketen kompatibel ist. Wie auch bereits beim Content, sollte die inkrementelle Erweiterung und Aktualisierung der Software innerhalb des personalisierten Systems realisiert werden. Dadurch wird ein ständiges Neuerstellen des Systems verhindert.

Mit wachsender Anzahl der Nutzer und der damit verbundenen Meta-Daten können auch die Funktionsweise der Content-Shops und des Empfehlungssystems angepasst werden.

Sowohl die zur Auswahl herangezogenen Algorithmen selbst als auch die Art des Systems können angepasst und überarbeitet werden. Hierfür kann es ggf. nötig sein, das Datenmodell zur Haltung der Metadaten anzupassen.

Durch Erweiterung der verwertbaren Daten um z.B. Prüfungsergebnisse oder die Wahl der Veranstaltungen, kann auch die Vielfalt der angebotenen Inhalte erhöht werden.

Die stetig wachsende Zahl an mobilen Endgeräten macht ebenfalls die Betrachtung einer möglichen Nutzung des bereitgestellten Systems auf unterschiedlichen mobilen Plattformen notwendig.

Nicht alle mobile Endgeräte bieten die Möglichkeit, OS-Images zu verwenden. Eine Möglichkeit, trotz dieser Einschränkung das erstellte OS-Image zu verwenden, ergibt sich dann, wenn die Images direkt von Rechenzentrum gehostet werden, und die Studierenden z.B. über eine Web-Desktop Lösung oder eine View-Umgebung auf Ihre personalisierten Systeme zugreifen können. Dies setzt zwar eine Internetanbindung voraus – spätestens seit der flächendeckenden W-LAN Anbindung auf dem Campus stellt dies jedoch kaum ein Problem dar.

## Literaturverzeichnis

- [HKR08b] T. Hildmann, O. Kao, and C. Ritter. Rollenbasierte Identitäts- und Autorisierungsverwaltung an der TU Berlin. 1. DFN-Forum Kommunikationstechnologien Verteilte Systeme im Wissenschaftsbereich, 2008.
- [HR07] T. Hildmann and C. Ritter. TUBIS-Integration von Campusdiensten an der Technischen Universität Berlin. *PIK-Praxis der Informationsverarbeitung und Kommunikation*, 30(3):145–151, 2007.
- [RHK10] C. Ritter, T. Hildmann, O. Kao. Erfahrungen und Perspektiven eines rollenbasierten IdM , 3. DFN-Forum Kommunikationstechnologie, 26. Mai 2010

# Anonymisierung und externe Speicherung in Cloud-Speichersystemen

Konrad Meier, Steffen Philipp

Albert-Ludwigs-Universität Freiburg  
Professur für Kommunikationssysteme  
Hermann-Herder-Str. 10  
79104 Freiburg  
konrad.meier@rz.uni-freiburg.de  
steffenmatthiasphilipp@googlemail.com

**Abstract:** Cloud-Speichersysteme ermöglichen es, Daten kostengünstig zu speichern und ortsunabhängig auf die gespeicherten Daten zuzugreifen. Einer Nutzung solcher Systeme durch die öffentliche Verwaltung und durch private Unternehmen stehen bisher oft datenschutzrechtliche Regelungen entgegen. Die vorliegende Arbeit analysiert die juristischen Anforderungen bei der Verwendung von Cloud-Speichersystemen und beschreibt ein reversibles Anonymisierungsverfahren, das es ermöglicht, personenbezogene Daten in Cloud-Speichersystemen abzulegen. Die datenschutzrechtlichen Regelungen werden dabei nicht verletzt. Das Verfahren wurde in einem föderierten Speichersystem implementiert und evaluiert.

## 1 Einleitung

In den vergangenen Jahren haben es Cloud-basierte Speichersysteme ermöglicht, die stark wachsende Datenmenge zu speichern und flexibel zur Verarbeitung bereitzustellen. Dabei skalieren die Systeme sowohl bei den Datenmengen als auch bei der Anzahl der Nutzer. Um dem kurzfristigen Speicherbedarf von Forschungs- und Projektgruppen gerecht zu werden, ist es notwendig, projektbezogene Speichercontainer flexibel bereitstellen zu können. Das Speichern und Auslagern von Daten in externen Cloud-Speichersystemen ist somit interessant, um lokale, stark genutzte Ressourcen kosteneffizient zu entlasten. Bedarfsspitzen einzelner Nutzer können mit geringem Aufwand abgefangen werden, was zu Kosteneinsparungen führt. Wesentliche Hindernisse, die einer externen Speicherung in solchen Systemen bisher entgegenstehen, sind neben allgemeinen Sicherheitsbedenken vor allem datenschutzrechtliche Regelungen, die für personenbezogene Daten gelten. Die Vorteile dieser Systeme führen aber oft dazu, dass Benutzer datenschutzrechtliche Aspekte ignorieren oder bewusst eine Verletzung dieser in Kauf nehmen<sup>1</sup>. Diese Problematik zeigt sich auch im Hochschul Umfeld.

Speziell für personenbezogene Daten müssen strenge datenschutzrechtliche Aspekte be-

---

<sup>1</sup> Beispiel: Verwenden von Dropbox für personenbezogene Daten.



achtet werden. Von zentraler Bedeutung ist hier die Regelung des § 4 Abs. 1 Bundesdatenschutzgesetz (BDSG), nach der jede Verarbeitung personenbezogener Daten, also auch deren Auslagerung in ein Cloud-System, entweder einer Erlaubnis durch Rechtsvorschrift oder einer wirksamen Einwilligung des Betroffenen bedarf. Da zum Zeitpunkt der Speicherung der Daten eine solche Einwilligung meist nicht gegeben ist und ein nachträgliches Einholen zu aufwändig ist, erscheint eine Auslagerung personenbezogener Daten in externe Cloud-Speichersysteme zunächst als nicht möglich. Für Daten ohne Personenbezug oder anonymisierte Daten trifft diese Einschränkung nicht zu. Die Praxis zeigt jedoch, dass eine Unterscheidung zwischen personenbezogenen Daten und Daten ohne Personenbezug oft nur schwer möglich ist. Gerade bei Benutzerdaten, beispielsweise Home-Laufwerke von Universitätsmitarbeitern und Studenten, ist eine Klassifikation der Daten meist unmöglich. In den meisten Fällen kann nur der Besitzer selbst angeben, ob personenbezogene Daten vorliegen.

Aufgrund dieser problematischen Ausgangslage beschreibt und evaluiert diese Arbeit ein auf Verschlüsselungs- und Fragmentierungstechniken basierendes Verfahren, mit dem Daten vor der externen Speicherung so umgewandelt werden, dass die externe Speicherung reversibel anonymisiert erfolgt. Somit fallen die extern gespeicherten Daten aufgrund der Anonymisierung nicht in den Anwendungsbereich geltender Datenschutzregelungen. Eine Auslagerung und Speicherung von Daten in externe Cloud-Speichersysteme wird somit auch für personenbezogene Daten ermöglicht.

Die Arbeit ist wie folgt strukturiert: Im Abschnitt 2 werden die Erkenntnisse verwandter Arbeiten untersucht. Im Abschnitt 3 erfolgt eine Analyse der rechtlichen Anforderungen, die bei der Speicherung in externen Systemen beachtet werden müssen. Im Abschnitt 4 wird die technische Umsetzung einer reversiblen Anonymisierung vorgestellt und im Abschnitt 5, im Rahmen eines föderierten Speichersystems, umgesetzt. Die Erkenntnisse der Arbeit werden im Abschnitt 6 zusammengefasst.

## **2 Verwandte Arbeiten**

Cloud-Speichersysteme werden von vielen kommerziellen Anbietern als Dienste angeboten. Dies hat dazu geführt, dass zahlreiche Arbeiten zu den Themen vendor lock-in, Verfügbarkeit und Datenschutz veröffentlicht worden sind, um problematische Eigenschaften dieser Systeme zu beheben.

Eine Gegenüberstellung unterschiedlicher Arbeiten zum Thema verteilte Cloud-Speichersysteme sowie dem Cloud of Clouds-Ansatz ist die Arbeit von Slamanić und Hanser [SH12]. Dabei werden die betrachteten Arbeiten unter anderem hinsichtlich ihrer Eigenschaften Vertraulichkeit, Integrität, Verfügbarkeit sowie Zugangskontrolle analysiert. Die Autoren kommen zu dem Ergebnis, dass Datenschutzaspekte zwar prinzipiell als wichtige Eigenschaften angesehen werden, jedoch wird dieser Aspekt in den meisten Arbeiten nicht ausreichend berücksichtigt. Insbesondere der Aspekt der access privacy wird in keiner der Arbeiten beachtet.

In der Arbeit von Sheng et al [SMGL11] wird eine Bit-interleaving-Technik verwendet,

um Daten aufzuteilen und um sie dann anschließend bei Cloud-Anbietern zu speichern. Es wird argumentiert, dass alleine durch das Aufteilen der Daten ein ausreichender Datenschutz gewährleistet werden kann, da ein möglicher Angreifer keinen Zugriff auf alle Daten besitzt. Einen ähnlichen Ansatz verfolgt auch die Arbeit von Abu-Lidbeh et al. [ALPW10]. Hier werden RAID-Techniken verwendet, um das Problem des vendor lock-ins zu verhindern. Dabei werden die Daten so auf mehrere Anbieter aufgeteilt, dass auch der Ausfall eines Anbieters die Verfügbarkeit der Daten nicht gefährdet.

Der Cloud of Clouds-Ansatz wird auch in der Arbeit von Bessani et al. [BCQ<sup>+</sup>11] beschrieben. Auch hier werden die Daten auf unterschiedliche Cloud-Anbieter verteilt. Die Daten werden jedoch zusätzlich verschlüsselt. Das System soll die Verfügbarkeit, Vertraulichkeit und Verfügbarkeit der abgelegten Daten verbessern.

Ein dezentrales Speichersystem, basierend auf Erasure Codes, wird in der Arbeit von Yao et al. [YXH13] beschrieben. Das System soll dabei die Benutzerdaten über Verschlüsselung schützen. Sowohl die Daten als auch der Schlüssel werden dezentral gespeichert, um zu verhindern, dass ein nicht vertrauenswürdiger Speicherserver die Daten entwendet.

Alle zuvor aufgeführten Arbeiten bieten zwar einen Ansatz für ein dezentrales Speichersystem. Für den konkreten Anwendungsfall der personenbezogenen Daten ist jedoch eine juristische Analyse der Datenschutzproblematik für Deutschland notwendig. Nur so kann sichergestellt werden, dass eine technischen Lösungen auch gesetzlichen Regelungen genügt. Die vorliegende Arbeit bietet genau dies, indem sie, basierend auf einer juristischen Anforderungsanalyse, eine technischen Lösung präsentiert.

### **3 Rechtliche Anforderungen**

Im Zusammenhang mit der Datenverarbeitung in Cloud-Systemen ist eine ganze Reihe datenschutzrechtlicher Vorschriften zu beachten. Insbesondere sind dies Gesetzesvorschriften im Telekommunikationsgesetz (TKG), im Telemediengesetz (TMG) und im Bundesdatenschutzgesetz (BDSG). Speziell bei der Auslagerung personenbezogener Daten in Cloud-Speichersysteme sind jedoch in der Regel nur die Vorschriften des BDSG anwendbar<sup>2</sup>. Im Folgenden wird die datenschutzrechtliche Problematik einer Auslagerung personenbezogener Daten in öffentliche Cloud-Speichersysteme anhand der Regelungen des BDSG erläutert.

#### **3.1 Auftragsdatenverarbeitung**

Das Auslagern personenbezogener Daten ist prinzipiell im Rahmen einer Auftragsdatenverarbeitung nach § 11 BDSG denkbar. Im Fall einer Auftragsdatenverarbeitung wird ein externer Auftragnehmer (Cloud-Anbieter) nicht als Dritter, sondern als verlängerter Arm

---

<sup>2</sup> Heidrich und Wegener [HW10], S. 803, (806)

der intern verantwortlichen Stelle behandelt, § 11 Abs. 1 S. 1 BDSG und § 3 Abs. 8 S. 3 BDSG. Das BDSG stellt für eine wirksame Auftragsdatenverarbeitung strenge Anforderungen hinsichtlich der Auswahl des Auftragnehmers und der vertraglich zu treffenden Regelungen, § 11 Abs. 2 S. 1 und S. 2 BDSG. Dadurch wird sichergestellt, dass der Auftragnehmer sich tatsächlich wie der verlängerte Arm der verantwortlichen Stelle verhält. Hieraus ergeben sich jedoch Pflichten für den Auftraggeber, die bei Cloud-Speicher-Anbietern nur schwer eingehalten werden können. So müssen Kontroll- und Dokumentationsvorschriften eingehalten werden. Bei Cloud-Speichersystemen besteht oft schon das Problem, den Speicherort zu bestimmen. Die Kontrolle von technischen und organisatorischen Maßnahmen zum Schutz der Daten stellt den Auftraggeber zusätzlich vor das Problem, dass er keinen Zugriff auf die Verarbeitungsprozesse und Verarbeitungsanlagen des Cloud-Anbieters hat.

Die mit der Auslagerung von Daten in öffentliche Cloud-Speichersysteme angestrebte Kostenersparnis ist nur realisierbar, wenn die Auslagerung nicht in den Anwendungsbereich der genannten Datenschutzvorschriften fällt. Der zusätzliche Aufwand für eine Auftragsdatenverarbeitung steht sonst im Widerspruch zur gewünschten Kostenersparnis. Damit die genannten Datenschutzvorschriften nicht zutreffen, muss sichergestellt werden, dass es sich bei den auszulagernden Daten nicht (mehr) um personenbezogene Daten handelt. Hierfür muss jedoch zunächst geklärt werden, wie personenbezogene Daten definiert werden.

### **3.2 Personenbezogene Daten**

Zu den personenbezogenen Daten zählen Informationen, die einer natürlichen Person zugeordnet sind oder leicht einer Person zugeordnet werden können, wie beispielsweise eine Wohnadresse, eine Telefonnummer oder eine den Inhaber benennende E-Mail-Adresse<sup>3</sup>. Im Einzelnen ist oft schwer einzuschätzen, ob personenbezogene Daten vorliegen oder nicht. Auch der Begriff der personenbezogenen Daten selbst ist (wohl auch aufgrund der erheblichen Rechtsfolgen, die damit verknüpft sind) sehr umstritten. Das BDSG definiert personenbezogene Daten in § 3 Abs. 1 wie folgt:

Personenbezogene Daten sind Einzelangaben über persönliche oder sachliche Verhältnisse einer bestimmten oder bestimmbaren natürlichen Person (Betroffener).

Jedoch ist im BDSG nicht explizit geregelt, auf welche Stelle und auf welche Mittel es für das Kriterium „bestimmbar“ ankommen soll. Dazu gibt es zwei prinzipiell verschiedene Ansichten.

Nach der „relativen“ Ansicht sind für das Kriterium „bestimmbar“ nur die verantwortliche Stelle und deren Mittel relevant. „Verantwortliche Stelle“ ist jede Person oder Stelle, die personenbezogene Daten für sich selbst erhebt, verarbeitet oder nutzt oder dies durch andere im Auftrag vornehmen lässt, § 3 Abs. 7 BDSG<sup>4</sup>.

<sup>3</sup> Siehe BDSG § 3 Abs. 1

<sup>4</sup> Siehe hierzu auch Gola u. a. [GKK12], § 3 Rn. 10

Nach der „objektiven“ Ansicht sind für das Kriterium „bestimmbar“ auch Dritte und deren Mittel relevant. Ein Dritter in diesem Sinne ist außer dem Betroffenen jede Person oder Stelle außerhalb der verantwortlichen Stelle, § 3 Abs. 8 BDSG<sup>5</sup>.

Wird jedoch der Text der europäischen Datenschutzrichtlinie 95/46/EG [EU-95] hinzugenommen, kommt man zu dem Schluss, dass weder die „relative“ noch die „objektive“ Ansicht mit der Richtlinie 95/46/EG und dem BDSG vereinbar ist. Die systematische und richtlinienkonforme Auslegung führt vielmehr zu dem Ergebnis, dass es für das Kriterium „bestimmbar“ im Sinne des § 3 Abs. 1 BDSG darauf ankommt, ob die verantwortliche Stelle eine Person direkt oder indirekt, mit Mitteln der verantwortlichen Stelle oder mit Mitteln eines Dritten identifizieren kann, wobei alle Mittel zu berücksichtigen sind, die vernünftigerweise eingesetzt werden könnten.

Welche Mittel „vernünftigerweise“ eingesetzt werden könnten, wird mittelbar durch § 3 Abs. 6 BDSG konkretisiert. Dort werden, als Gegenstück zu personenbezogenen Daten, anonymisierte Daten definiert.

### 3.3 Anonymisierte Daten

Personenbezogene Daten sind Einzelangaben über persönliche oder sachliche Verhältnisse einer bestimmten oder bestimmbaren natürlichen Person. Bestimmbar ist eine Person, wenn die verantwortliche Stelle direkt oder indirekt, mit Mitteln der verantwortlichen Stelle oder mit Mitteln eines Dritten, eine Person identifizieren kann und der Aufwand für den Einsatz dieser Mittel im Hinblick auf Zeit, Kosten und Arbeitskraft nicht unverhältnismäßig ist, § 3 Abs. 6 BDSG.

Um Daten im Sinne des Datenschutzgesetzes zu anonymisieren, muss der Aufwand zum Identifizieren der Person unverhältnismäßig hoch sein. Was als unverhältnismäßig angesehen wird, ist im Gesetz nicht definiert. Eine Verschlüsselung personenbezogener Daten kann in diesem Zusammenhang nicht als Anonymisierung betrachtet werden<sup>6</sup>, da dies hauptsächlich einen unverhältnismäßig hohen Rechenaufwand (Aufwand an Zeit und Kosten) bedeutet. Ein unverhältnismäßig hoher Aufwand an Arbeitskraft ist nicht gegeben. Auch ist es möglich, dass zu einem späteren Zeitpunkt der hohe Zeitaufwand drastisch reduziert werden kann, wenn Schwachstellen im verwendeten Verschlüsselungsverfahren entdeckt werden<sup>7</sup>. Ein Angreifer könnte somit Daten entwenden und zu einem späteren Zeitpunkt das Verschlüsselungsverfahren brechen, wenn dieses über technische Mittel möglich geworden ist.

Ein technischer Lösungsansatz zur automatischen Anonymisierung kann darin bestehen, die Verschlüsselung durch Maßnahmen zu ergänzen, die sicherstellen, dass für einen Angreifer ein unverhältnismäßiger Aufwand an Arbeitskraft erforderlich ist und der erforder-

<sup>5</sup> Siehe hierzu auch Weichert in Däubler u. a. [DWWK10], § 3 Rn. 13

<sup>6</sup> Zu diesem Ergebnis kommt auch das Beratungsgremium, das von der europäischen Kommission auf Grundlage der Richtlinie 95/46/EG eingesetzt wurde, in [AA07a], S. 15, allerdings ohne nähere Begründung. Differenzierter war die Einschätzung des Gremiums noch in [AA07b], S. 24, wo abhängig von den Umständen noch von einer Anonymisierung ausgegangen wurde.

<sup>7</sup> So auch Spies [Spi11], Ziffer 2

liche Aufwand an Zeit und Kosten selbst dann unverhältnismäßig bleibt, wenn die Entzifferung für sich genommen nicht mehr mit unverhältnismäßigem Rechenaufwand verbunden ist. Nur so kann ein unverhältnismäßig hoher Aufwand für die drei Aspekte Zeit, Kosten und Arbeitskraft hergestellt werden. Dies ermöglicht eine gesetzeskonforme reversible Anonymisierung mit technischen Mitteln.

## **4 Technische Lösung: Reversible Anonymisierung**

Eine reversible Anonymisierung und damit die Erfüllung datenschutzrechtlicher Vorgaben, kann erreicht werden, wenn Datensätze nach dem Stand der Technik verschlüsselt, in mehrere Fragmente zerlegt und die Fragmente anschließend in voneinander unabhängigen Speichersystemen mit unabhängigen Zugriffssicherungen nach dem Stand der Technik gespeichert werden. Die Fragmentierung kann zufällig bitweise oder blockweise, beispielsweise byteweise erfolgen.

Durch die Speicherung der Fragmente in unabhängigen Speichersystemen entsteht eine zusätzliche, in Beschaffungsaufwand bestehende Hürde, die sich qualitativ von den Sicherheitskomponenten Verschlüsselung und Fragmentierung unterscheidet. Sie stellt sicher, dass keine extern speichernde Stelle die Originaldaten alleine mit Rechenaufwand wiederherstellen kann. Die so hinzugefügte Sicherheitskomponente ist nicht durch eine Schwächung der verwendeten Verschlüsselungs- und Fragmentierungsverfahren betroffen und stellt in Kombination mit Verschlüsselung und Fragmentierung sicher, dass die Herstellung des Personenbezugs für jede andere als die auslagernde Stelle im Hinblick auf Zeit, Kosten und Arbeitskraft auf Dauer unverhältnismäßig ist.

Ein Angreifer müsste somit die beteiligten externen Speichersysteme identifizieren, deren Zugriffssicherungen überwinden, alle Fragmente kopieren, die Fragmente zusammensetzen und die Verschlüsselung brechen.

Werden Datensätze wie beschrieben verschlüsselt, fragmentiert und gespeichert, sind die Daten in den Fragmenten anonymisiert. Die Übermittlung der Fragmente an externe Speicher-Systeme und deren Speicherung in externen Speichersystemen fällt nicht mehr in den Anwendungsbereich des BDSG.

Fragmente, die nach dem zuvor beschriebenen Verfahren erzeugt wurden, können, ohne Beachtung der für personenbezogene Daten einzuhaltenden rechtlichen Vorgaben, an externe Speichersysteme übermittelt und dort gespeichert werden. Die intern verantwortliche Stelle muss die Vorschriften des BDSG jedoch weiterhin beachten. Denn die ausgelagerten Daten sind, wie zuvor dargelegt, nicht für die intern verantwortliche Stelle anonymisiert, sondern nur für jede andere Stelle. Bei der Auswahl der externen Speicheranbieter sollte beachtet werden, dass eine Kooperation der Anbieter unwahrscheinlich ist.

Weitere Anforderungen ergeben sich aus der Regelung in § 30 Abs. 1 BDSG, die schon im Zusammenhang mit der Bestimmung des Begriffs der personenbezogenen Daten relevant geworden war. Sie legt fest, dass die für eine Identifizierung erforderlichen Zusatzinformationen separat zu speichern sind. Dies bedeutet im Fall der reversiblen Anonymisierung, dass die Informationen zum Wiederherstellen der Daten nicht zusammen mit den Frag-

menten ausgelagert werden dürfen. Daher werden diese Daten im lokalen Speichersystem abgelegt.

## 5 Föderiertes Speichersystem

Das föderierte Speichersystem verbindet mehrere Cloud-Speichersysteme und deren jeweilige Systeme zu einem Speicherverbund [MWS13]. Mit diesem System ist es beispielsweise möglich, die Speicherressourcen an mehreren Universitäten im Verbund zu nutzen. Das System basiert auf dem Ansatz eines Object-Storage-Systems mit frei wählbarem Datenspeicherort und wird durch eine Abstraktions- und Verwaltungsschicht über der lokalen Speicherverwaltung ermöglicht. Die Abstraktionsschichten und die Kommunikation zwischen ihnen ist in Abbildung 1 dargestellt. Die Abbildung zeigt für die Rechenzentren A und B das vorgeschlagene Schichtenmodell. Die Kommunikation zwischen den Rechenzentren erfolgt über die Schicht der föderierten Speicherverwaltung. Die Rechenzentren C und D sind äquivalent aufgebaut. Wie in der Abbildung zu sehen ist, setzt die föderierte

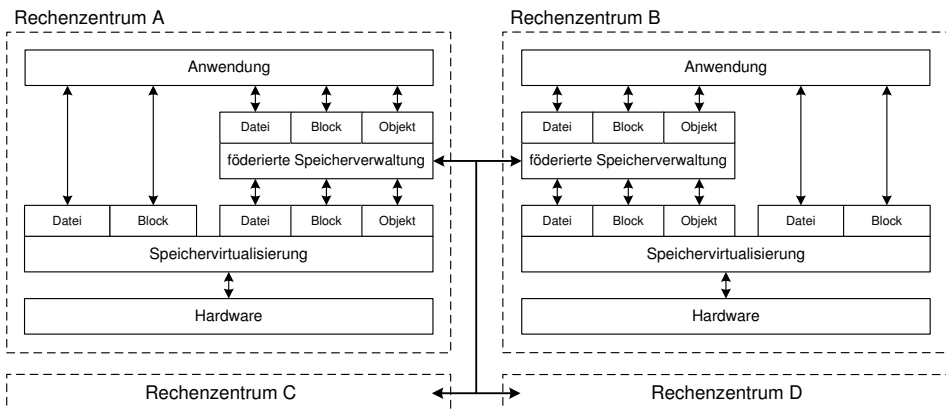


Abbildung 1: Schichtenmodell des föderierten Speichersystems mit Kommunikationsbeziehungen

auf eine lokale Speicherverwaltung auf und nutzt die bereitgestellten Schnittstellen, um auf Speichercontainer zuzugreifen. Applikationen greifen über die föderierte Speicherverwaltung auf die Speichercontainer zu. Eine detaillierte Beschreibung der einzelnen Schichten ist in der früheren Arbeit [MWS13] zu finden.

Jede Anwendung ist in der Lage, mit Hilfe von Metadaten zu definieren, welche Anforderungen sie an den Datenschutz ihrer gespeicherten Daten hat. Die föderierte Speicherverwaltung stellt anschließend sicher, dass die Datenschutzdefinition eingehalten wird, auch wenn die Daten den lokalen Standort verlassen und an einen weiteren Standort innerhalb der Föderation ausgelagert werden. Die Datenschutzdefinition umfasst aktuell die Möglichkeit, die Daten nach mehreren Sicherheitsleveln zu klassifizieren. Personenbezogene Daten werden dabei mit dem höchsten Sicherheitslevel abgespeichert.

Die Definition sieht aktuell wie folgt aus:

```
SecurityLevel = 0|1|2|9
0: Daten nicht verschlüsseln
1: Daten beim Auslagern verschlüsseln
2: Daten lokal und beim Auslagern verschlüsseln
9: personenbezogene Daten
```

Der in Kapitel 4 beschriebene Ansatz zur reversible Anonymisierung wurde im Rahmen dieses föderierten Speichersystems beispielhaft implementiert. Das föderierte Speichersystem basiert auf dem Object-Storage System OpenStack Swift<sup>8</sup>. In Swift-Systemen werden Objekte in Containern gespeichert und Container in Accounts. Objekte umfassen Daten (Content) und Metadaten (Header). Die Metadaten werden unter anderem verwendet, um die definierten Sicherheitslevel zu speichern. Eine Sicherheitslevel-Definition wird als String gespeichert und kann wie folgt aussehen:

```
X-Object-Meta-Security = "9"
```

Eine strukturelle Übersicht der Komponenten der föderierten Speicherverwaltung ist in Abbildung 2 gegeben. Wie der Grafik entnommen werden kann, dient der Federated-Proxy

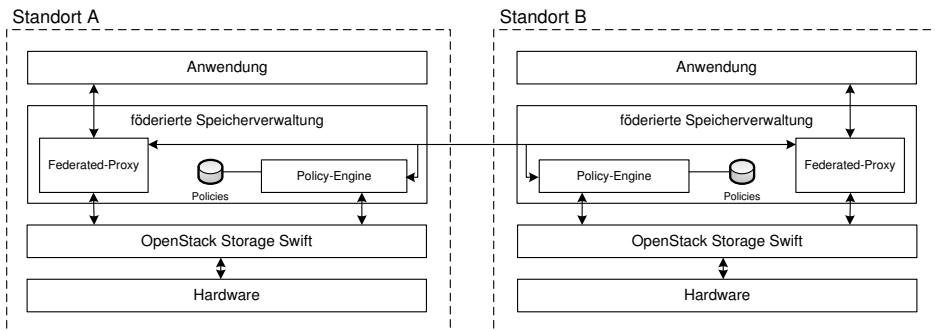


Abbildung 2: Komponenten der föderierten Speicherverwaltung

den Applikationen als Schnittstelle zum föderierten Speichersystem. Er sorgt dafür, dass Daten entsprechend ihrer Sicherheitsstufe abgelegt werden.

Das Policy-Tool liest Statistik-Daten aus dem Speichersystem aus und bestimmt anhand hinterlegter Policies immer wieder neu, in welchem Speichersystem Daten abzulegen sind. Policies sind automatische Regeln, die es ermöglichen, Daten zwischen den Standorten zu verschieben. Dies ermöglicht es, Daten, die nicht länger lokal benötigt werden, in einen Überlaufspeicher auszulagern. Das Verschieben der Daten an einen anderen Standort wird als "export" bezeichnet. Die entsprechende Umkehrfunktion wird als "import" bezeichnet.

<sup>8</sup> Object Storage System OpenStack Swift: <http://swift.openstack.org>

Der Syntax der Policies ist wie folgt definiert:

|                                                                                                                                           |                                                                                                         |
|-------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------|
| <pre>RULE ruleName   EXPORT [source storage condition]   FROM containerName [data condition]   TO target [target storage condition]</pre> | <pre>RULE ruleName   IMPORT [source storage condition]   CONTAINER containerName [data condition]</pre> |
|-------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------|

Die Bedingungen in eckigen Klammern sind optional. Der Name der Policies, `ruleName`, ist notwendig, um die Regeln zu unterscheiden. Danach wird die Aktion angegeben und gegebenenfalls mit einer Bedingung verknüpft (beispielsweise mehr als 80 % der lokalen Speicher-Ressourcen ist aufgebraucht). Die Daten, die verschoben werden sollen, werden als `containerName` angegeben und können über die `data condition` eingeschränkt werden. So ist es beispielsweise möglich, nur Daten zu verschieben deren Änderungsdatum älter ist als ein bestimmtes Datum. Zuletzt wird das Ziel-Speichersystem `target` angegeben. Auch hier ist es wiederum möglich, eine Bedingung anzugeben (z. B. Größe des noch verfügbarer Speichers).

Der Transporter verlagert Daten entsprechend ihrem Sicherheitslevel und in Übereinstimmung mit datenschutzrechtlichen Anforderungen in und zwischen den Swift-Systemen. Er verfügt über einen Load-Balancer und mehrere Worker. Die Worker sind jeweils als separate Prozesse mit mehreren Threads implementiert und kommunizieren über Sockets mit dem Load-Balancer. Der Transporter implementiert das Auslagern von personenbezogenen Daten über die reversible Anonymisierung.

Werden Daten durch eine Applikation im föderierten Speichersystem abgelegt, geschieht dies über den Federated-Proxy. Dieser behandelt alle Daten entsprechend ihrer Metadaten. Wird kein Sicherheitslevel angegeben, behandelt der Federated-Proxy die Daten wie personenbezogene Daten.

1. Datensatz empfangen
2. Verschlüsseln, wenn `SecurityLevel = 2` oder `9`
3. Datensatz mit Metadaten lokal schreiben

Falls verschlüsselt wird, wird die symmetrische Blockchiffre AES verwendet. Dabei wird der Objekt-Header um die zusätzlichen Metadaten für die Verschlüsselung erweitert. Die Schlüssel für die Verschlüsselung werden in der lokalen Benutzerverwaltung gespeichert.

## 5.1 Datenschutzkonformes Auslagern von Daten

Der Export von Daten in externe Speichersysteme innerhalb der Föderation wird über ein Transport-Kommando des Policy-Tools ausgelöst und vom Transporter ausgeführt. Das in Kapitel 4 beschriebene reversible Anonymisierungsverfahren wird angewendet, wenn personenbezogene Daten (`SecurityLevel 9`) an externe Speichersysteme weitergegeben werden. Dabei wird folgender Ablauf angewandt:

1. Verschlüsselten Datensatz mit Metadaten lokal lesen
2. Verschlüsselten Datensatz fragmentieren
3. Fragmente extern speichern
4. Link-Informationen verschlüsseln



## 5. Verschlüsselte Link-Informationen lokal speichern

Die Fragmentierung erfolgt im aktuellen Prototyp in 8 Bit-Blöcken. Die Verteilung der Blöcke auf die Fragmente wird über ein Fragmentierungsmuster aus Zufallszahlen bestimmt. Sollte zu einem Zeitpunkt die verwendete Verschlüsselungsmethode unsicher geworden sein, bleiben die in den Fragmenten gespeicherten Daten anonymisiert. Da keines der Fragmente die verschlüsselten Daten enthält, kann eine Entzifferung erst nach der Defragmentierung erfolgen. Die Anzahl der Fragmente wird aktuell statisch festgelegt und stellt somit eine harte Anforderung an die benötigte Anzahl externer Speichersysteme. Dabei wird genau ein Fragment in einem externen Speichersystem abgelegt. Stehen eine unzureichende Zahl externer Speichersysteme zur Verfügung, können die Daten nicht ausgelagert werden. Bei der aktuell verwendeten Fragmentierungstechnik werden die Daten ohne Redundanz aufgeteilt und exportiert. Beim Zugriff auf die Daten müssen alle externen Speichersysteme verfügbar sein. Um diese Beschränkung zu umgehen, können RAID-ähnliche Fragmentierungstechniken eingesetzt werden, die über Redundanzen den Ausfall kompensieren können.

Um die Daten später wieder zusammensetzen zu können, werden Link-Informationen im lokalen Speichersystem abgelegt. Diese werden im Datenbereich (Content) des ursprünglichen Objekts gespeichert. Sie enthalten Metadaten, die auf den Speicherort, den Container und den Objektnamen der Fragmente im externen Speichersystem verweisen. Der externe Objektname wird über einen String gebildet, der aus dem internen Speicherzeitpunkt, der Kopie-Nummer, dem internen Objektnamen und der Fragment-Nummer besteht. Dadurch wird für externe Angreifer die Zuordnung der Fragmente zueinander erschwert.

### 5.2 Lesen von Daten

Der Zugriff auf die Daten erfolgt immer über den Federated-Proxy. Dieser liest die Metadaten des angeforderten Objekts, um zu unterscheiden, ob die Daten lokal oder extern abgelegt sind.

Sind die Daten lokal gespeichert, ist es möglich, dass die Daten verschlüsselt vorliegen (SecurityLevel 2 oder 9). In diesem Fall werden die Daten entschlüsselt und der Anwendung bereitgestellt.

Sind die Daten in ein externes Speichersystem ausgelagert worden, werden die lokalen Link-Informationen gelesen. Anschließend werden die Fragmente aus den externen Systemen gelesen, defragmentiert, entschlüsselt und an die Anwendung übertragen. Der Federated-Proxy stellt insgesamt sicher, dass der Nutzer bzw. die Applikation die zum Speichern an den Federated-Proxy übermittelten Daten in dieser Form auch zurückerhält.

## 6 Fazit und Ausblick

Der Verwendung von Cloud-Speichersystemen standen bisher meistens Datenschutzbedenken gegenüber. Diese Bedenken sind gerade bei personenbezogenen Daten absolut gerechtfertigt und lassen sich juristisch begründen. Da selbst eine Verschlüsselung personenbezogener Daten nicht als ausreichender Schutz beim Auslagern von Daten angesehen werden kann, schien bisher eine Verwendung solcher Systeme nur im Rahmen einer Auftragsdatenverarbeitung möglich.

Diese Arbeit zeigt einen alternativen Ansatz, um eine problematische Auftragsdatenverarbeitung zu umgehen. Das vorgestellte Verfahren der reversiblen Anonymisierung beschreibt, wie Daten derart verändert werden können, damit sie nach den Anforderungen des BDSG als anonymisiert betrachtet werden können. Die Daten werden dabei nach dem Stand der Technik verschlüsselt, in mehrere Fragmente zerlegt und die Fragmente anschließend in voneinander unabhängigen Speichersystemen mit unabhängigen Zugriffssicherungen nach dem Stand der Technik gespeichert. Die beispielhafte Umsetzung dieses Verfahrens in einem föderierten Speichersystem hat gezeigt, dass eine solche Methode praktisch umgesetzt werden kann. Das föderierte Speichersystem übernimmt die Aufgabe, die Daten entsprechend ihrer Klassifizierung zu schützen, selbst wenn die Daten den ursprünglichen Standort verlassen.

## Literatur

- [AA07a] Artikel-29-Arbeitsgruppe. Opinion 05/2012 on Cloud Computing, WP 196. *Drucksachen Europäische Kommission*, 2007.
- [AA07b] Artikel-29-Arbeitsgruppe. Stellungnahme 4/2007 zum Begriff „personenbezogene Daten“, WP 136. *Drucksachen Europäische Kommission*, 2007.
- [ALPW10] Hussam Abu-Libdeh, Lonnie Princehouse und Hakim Weatherspoon. RACS: A Case for Cloud Storage Diversity. In *Proceedings of the 1st ACM Symposium on Cloud Computing*, SoCC '10, Seiten 229–240, New York, NY, USA, 2010. ACM.
- [BCQ<sup>+</sup>11] Alysson Bessani, Miguel Correia, Bruno Quaresma, Fernando André und Paulo Sousa. DepSky: dependable and secure storage in a cloud-of-clouds. In *Proceedings of the sixth conference on Computer systems*, EuroSys '11, Seiten 31–46, New York, NY, USA, 2011. ACM.
- [DWWK10] Wolfgang Däubler, Peter Wedde, Thilo Weichert und Thomas Klebe. *Bundesdatenschutzgesetz - Kompaktkommentar zum BDSG und anderen Gesetzen*. Bund-Verlag, 2010.
- [EU-95] EU-Parlament. Richtlinie 95/46/EG. *Amtsblatt*, (L 281):31–50, Oktober 1995. Richtlinie 95/46/EG des Europäischen Parlaments und des Rates vom 24. Oktober 1995 zum Schutz natürlicher Personen bei der Verarbeitung personenbezogener Daten und zum freien Datenverkehr.
- [GKKS12] Peter Gola, Christoph Klug, Barbara Körffer und Rudolf Schomerus. *BDSG Bundesdatenschutzgesetz*. Gelbe Erläuterungsbücher. Beck C. H., 2012.

- [HW10] Joerg Heidrich und Christoph Wegener. Sichere Datenwolken - Cloud Computing und Datenschutz. In *MultiMedia und Recht*, 2010.
- [MWS13] Konrad Meier, Dennis Wehrle und Nico Schlitter. Ein Konzept zum Aufbau eines föderierten, dezentralen Speichersystems im Hochschulumfeld. In *6. DFN-Forum Kommunikationstechnologien*. GI, 2013.
- [SH12] D. Slamanig und C. Hanser. On cloud storage and the cloud of clouds approach. In *International Conference for Internet Technology and Secured Transactions*, Seiten 649–655, 2012.
- [SMGL11] Zhonghua Sheng, Zhiqiang Ma, Lin Gu und Ang Li. A privacy-protecting file system on public cloud storage. In *International Conference on Cloud and Service Computing (CSC)*, Seiten 141–149, 2011.
- [Spi11] Axel Spies. Cloud Computing: Keine personenbezogenen Daten bei Verschlüsselung. *MMR-Aktuell*, 2011.
- [YXH13] Chuan Yao, Li Xu und Xinyi Huang. A Secure Cloud Storage System from Threshold Encryption. In *5th International Conference on Intelligent Networking and Collaborative Systems (INCoS)*, Seiten 541–545, 2013.

# Mapping virtual paths in virtualization environments

Martin G. Metzker, Dieter Kranzlmüller

Institut für Informatik  
Ludwig-Maximilians-Universität München  
Oettingenstraße 67  
80538 München  
{metzker, kranzlmueeller}@mnm-team.org

**Abstract:** Quality of Service (QoS) in networks is an essential ingredient for providing attractive, reliable services to customers. In contrast to model layered protocol stacks, where resources are used exclusively, virtualization introduces a series of novel specific challenges, due to shared resources. This paper presents the model and mapping procedure of an approach for matching network QoS attributes with the needs of virtualization environments (VE). The key of our solution is an extended view on network topologies which adequately includes virtual components and links. We identify discernible types of network components and links, to enable accurate descriptions of network paths through VEs. Our procedure is used to automatically describe (sub-)topologies yielding enough information to enable network QoS implementations.

## 1 Introduction

Over the last few years, virtualization with its plethora of applications has been changing the development of IT infrastructures in data centres. Within a virtualization environment (VE), *virtual infrastructures* (VI), consisting of *virtual machines* (VM) and virtual network components, are created, managed and deployed as needed for predefined services or simply as a vanilla topology of servers for arbitrary use through customers.

Specifications and management of VIs are abstracted from the underlying physical infrastructure and intentionally leave out details about the concrete implementation. When activated, VIs are placed within the VE, onto the providing physical components. This includes two important aspects:

1. Placing a VI within a VE usually expands its network links to network paths, including many physical and virtual components.
2. *Quality of Service* (QoS) attributes and requirements are formulated for abstract components and links, requiring refinement to match their placement.

Many aspects of VIs can be changed swiftly and on demand, especially the shape and participants of their respective virtual network topologies. The implied changes on resource consumption and consequently available network capacities, may lead to VMs starving each other for network resources. In extreme cases the services they provide are no longer

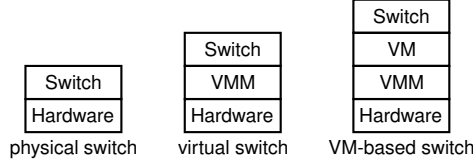


Figure 1: Three methods of implementing a switch

useful. In order to avoid such effects, network QoS management must be performed.

This contribution introduces a model and procedure to map VI network paths onto VEs. This enables automated mapping of network QoS requirements, making them an integral part of a management architecture for network QoS management in VEs. The procedure is intended to be automated, so that it can be applied for all VI paths. We prove our methodology by exemplary introducing the throughput QoS attribute to a Xen based VE.

## 2 Analysis and Approach

Looking at traditional network QoS with only physical components, network components and endpoints are physical resources with known capacities to provide services. In VIs, components, endpoints and links can be placed on dedicated physical hardware, or as virtual instances, implemented as a service on a shared platform.

Network virtualization yields a variety of network components, that are equivalent regarding their purpose and functionality according to the ISO OSI model, but may mandate different treatment for network QoS and QoS management. Figure 1 exemplary shows three conceptual implementations of a network switch. The physical switch, is implemented on dedicated hardware, whereas the virtual switch is one of possibly many virtual components realised by a *virtual machine monitor* (VMM), installed on a physical computer (*host*). VM-based switches are realised as part of a VM’s operating system.

Sharing few physical components among many virtualized entities involves a multiplex which must be managed, in order to correctly manage network QoS. Often, multiplexes are hidden and we attribute this challenge being unsolved to the increasing complexity introduced by virtual components and shared resources.

Performing network QoS management concerns different (sub-)components and therefore requires an adopted management procedure. For each example in Figure 1, each layer represents (at least) a subcomponent of the switch, possibly requiring additional management when management on the switch is performed. This is a challenge for planning and provisioning, as the managed component “switch” becomes one of many possible concrete implementations. Further, in VEs, single components may be moved to another host after initial deployment (*migration*), which may result in a change of the implementation of “switch”, and therefore how QoS must be implemented, enforced and monitored.

This gap between VI and component management can be met by requiring static configurations, eliminating a key feature of virtualization, or by a system with abstract description of VIs' QoS paths as *target configuration* that are continuously adopted to the VE's current state. We aim to realise the latter, for which we identify four main steps:

- |                                          |                                        |
|------------------------------------------|----------------------------------------|
| S1 determining involved components,      | S3 adopt to contestants and shortages, |
| S2 mapping of high level QoS attributes, | S4 develop monitoring strategies.      |

This contribution focuses on S1 to determine all involved components of a VI network path and their types, so that QoS management can be performed correctly, with respect to the VE's current state.

The overall goal is an adoption of the layered approach first introduced in [JN04], that develops “concrete QoS parameters [...] without] knowledge of the underlying [...] network conditions” as an intermediary result. In our adoption we aim to map specifications of QoS parameters for VIs onto the VE. Referring to the “layers [representing] very different features of services” in [JN04], our approach corresponds to developing mappings from the *application-layer* to the *resource-layer*.

We follow a bottom up approach and first analyse types of network components and links that differ in QoS management. The result is an enhanced topology resource-layer model, containing information on the degree of virtualization of components. This enables canonical implementation by the concrete technologies and components used to build VEs. The introduced procedure describes the refinement from application- to resource-layer.

## 2.1 Components

Figure 2 shows a classification of network components by locality of configuration, introduced in [MD10]. *Physical* components are fixed within the VE. While they may be replaced, for instance due to hardware failures, their place and role within the network topology remains unchanged. *Local* components, are placed on hosts. Their roles and function may be copied to other hosts but their state is relevant only for their current host. *Volatile* components are (equivalent to) VMs in terms of placement within the VE. Their state is invariant to placement within the VE, so they can be integrated at different positions in the (virtual) network topology.

This classification indicates how much immediate control a component has over the underlying physical resources. While physical components can map available resources to network performance accurately, local components contest for shared resources with other components and VMs. Volatile components are in general unaware of physical resources and how they are coordinated. Following through, the introduced classification separates how network QoS must be implemented by components:

- Physical components can directly implement network QoS.
- Local components must additionally handle resource contention. Effectively they implement network QoS with respect to the current load of their host.
- Volatile components implement some QoS themselves, but must delegate the task of ensuring that the required resources are available to the virtualization platform.

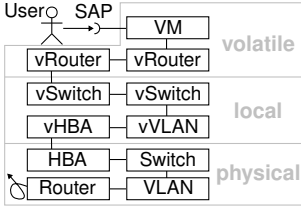


Figure 2: A conceivable path across network component classes and endpoints.

| src ↓                       | dst →       |               |               |
|-----------------------------|-------------|---------------|---------------|
|                             | p           | l             | v             |
| p                           | (a) p. link | (b) p. uplink | (d) pv-bridge |
| l                           | (b)         | (c) E2E link  | (e) v. link   |
| v                           | (d)         | (e)           | (f) RDMA      |
| physical   local   volatile |             |               |               |

Table 1: Generic discernible direct link types by end-point types

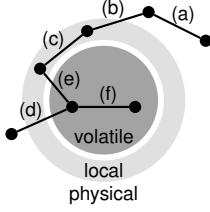
## 2.2 Links

As their physical role models, virtual network links are passive and network QoS management is performed at their endpoints. For any functional kind of network component, e.g. switch, classification by locality of configuration yields different types, concerning how resources are committed and network QoS management is performed. Linking components of different types, yields a set of discernible link types, summarised in Table 1.

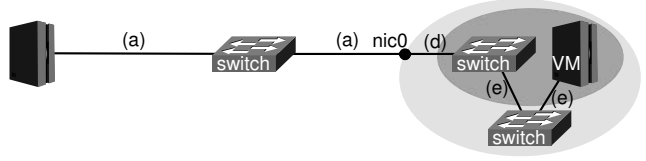
To characterise these six link types, we look at how the endpoints and therefore the link are managed. The types *physical link*, *physical uplink*, *E2E link* correspond to links also found in strictly physical infrastructures, while the types *pv-bridge*, *virtual link*, *RDMA* bear virtualization specific challenges. We characterize the link types as follows:

- (a) **physical link** There aren't any virtual components involved, "traditional" QoS [ea98].
- (b) **physical uplink** Virtual components are connected to the physical network. The VMM can assign a virtual component resources like an OS scheduler to a process. For QoS management the virtual component can be regarded as an endpoint in strictly physical topologies.
- (c) **E2E link** A trivial end-to-end (E2E) link in the context of physical networks, as both endpoints are located on the same host.
- (d) **pv-bridge** The uplink of a volatile component to the physical topology. The volatile component is generally not aware of resource contestants, but has been given an uncontested direct uplink, through the VMM. Even though the uplink is uncontested, the VMM must ensure the volatile component has enough other resources available.
- (e) **virtual link** The usual way of connecting virtual components, mostly a volatile component and a local component. The VMM manages both endpoints and can allocate the resources to meet QoS demands.
- (f) **RDMA** A theoretical direct connection between volatile components, without data ever crossing the VMMs domain. It is called remote direct memory access (RDMA) as this is most likely how it would be implemented, as both endpoints are placed on the same host and traffic between them is never handed to the VMM.

This constitutes a complete list of link types regarding the virtualization of endpoints, with which network paths through VEs can be described topologically. For (QoS) management



(a) Links from Table 1



(b) Enhanced topological view with virtualization host introspection

Figure 3: Resource layer views

links also differ in the OSI layers the endpoints implement and the concrete technology employed for virtualization. The above list is a generic template and specialized versions for the specific technologies employed in the endpoints must be derived. Technology and OSI layer characteristics for QoS handling yield a VE specific finite set of discernible link types, for which management functions must be available.

### 2.3 Resource-Layer topology view

At resource-layer all relevant information is described non-ambiguously [JN04]. We propose a topological view which groups network components by host and distinguishes between local and volatile components. This is consistent with describing links by their endpoints and components by their degree of abstraction from the physical hardware.

Figure 3(a) shows an abstract example featuring each link named in Table 1. By making the endpoints' domains (physical, local, volatile) explicit, the various link types are discernible without the labels. Figure 3(b) keeps the labelling for clarity purposes. The Figure itself shows an application example, connecting two servers over three switches, where each switch is one of the implementations introduced in Section 2. The physical switch, next to the physical host is correctly displayed as independent and the VM-based switch is displayed in a different domain as the local switch, yet within the same host.

Coloured/shaded areas group local and volatile components by host. All virtual components realised by the same physical host belong to the same area. With this illustration, volatile components are always contained within the local area. They are set apart optically, by using a different shade. The link between the two virtualized switches crosses domains, which means it must be a type-e link. The intersection of a link with the host boundary in Figure 3(b) is marked and explicitly names the used host's NIC, *nic0*. Which NIC is used determines which resources are used for the specific link, thus must be represented on the resource layer.

This enhanced topology view may include all identified component and link types. The increase in represented information meets the increased complexity and qualifies as resource layer for QoS layer partitioning following [JN04] as described in Section 2.



### 3 Refinement procedure

With the resource-layer established in Section 2.1, we can outline a generic procedure for mapping application-layer management information or tasks to the resource-layer. As an example we use the mapping of management information in the form of QoS attributes. This is achieved by iteratively refining application-layer links to eventually match the resource-layer topology. QoS attribute refinement is then performed based on reasoning on intermediary results. With each step the QoS attribute and its requirements are adopted to suit the sub-path which ultimately results in QoS requirements for each link at resource layer, which can then be implemented in a technology specific manner.

Network components are often characterized by the highest layer of the ISO OSI reference model they implement, e.g. routers implement layer 3. Combined with our taxonomy in Section 2.1 every component within a virtualization infrastructure can be characterized by the highest OSI layer it implements and its locality of configuration. The refinement procedure must therefore trigger two tasks: adoption of QoS attributes to the identified link segments and advancing the refinement process by detecting and resolving compound links. Our procedure develops the resource-layer topology (final topology) from the application-layer topology (initial topology), based on available management information. Algorithm 1 shows the core procedure `refinePath`, which uses two noteworthy subroutines: `link` (lines 8 and 22) and `adoptQoS` (lines 8, 12 and 22).

```

input :  $T_{res}$  as the list of all links in the final topology,  $P_{input}$  as endpoints of a path and the path's QoS attribute
output:  $P_{refined}$  as a list of links in the final topology associated with an associated QoS attribute

1  $P_{refined} \leftarrow \epsilon$ 
2  $e_0, e_1 \leftarrow P_{input}.endpoints$ 
3  $P_{res-layer} \leftarrow pathAtOsiLayer(e_0, e_1, e_0.layer)$ 
4 while  $P_{res-layer} \notin T_{res}$  do
5    $n \leftarrow neighborOf(e_0, P_{res-layer})$ 
6   if  $n \neq e_1$  then
7     if  $path(e_0, n) \in T_{res}$  then
8        $P_{refined}.append(link(e_0, n, adoptQoS(P_{input}, n)))$ 
9     else
10       $P_{sub} \leftarrow path(e_0, n)$ 
11       $P_{sub,qos} \leftarrow adoptQoS(P_{input}, n)$ 
12       $P_{refined}.append(elementsOf(recursion into subpath))$ 
13    end
14     $e_0 \leftarrow n$ 
15  end
16  else if  $n = e_1$  then
17     $e_0.layer \leftarrow e_0.layer - 1$ 
18  end
19   $P_{res-layer} \leftarrow pathAtOsiLayer(e_0, e_1, e_0.layer)$ 
20 end
21  $P_{refined}.append(link(e_0, e_1, adoptQoS(P_{res-layer}, e_0.layer)))$ 
22 return  $P_{refined}$ 

```

**Algorithm 1:** The `refinePath` procedure

The `refinePath` procedure uses the property of the OSI model, that every layer  $n$  uses the links and services provided by layer  $n-1$ , to break paths into sub-paths, starting at the layer of an input-endpoint, proceeding to lower layers. If a sub-path cannot be split up, it corresponds to a single resource-layer link, for which adequate management functions are available. When a path with a directly corresponding link in the resource-layer is found, the function `link` is called. Calling `link` is a request to apply the input-path's

QoS attribute to a resource-layer link. To actually apply the attribute, `link` uses topology information to identify the technologies used for the endpoints and the link type. With this information `link` can determine the correct management for the path segment.

The task of `adoptQoS` is provide an adopted QoS attribute to `link`, by tracking the performed path refinement. The concrete adoption procedure is specific to the QoS attribute. For instance, for a sensible implementation for an attribute “data rate”, `adoptQoS` would alter the attribute values to account for protocol header fields, when regarding data rate at varying OSI layers, or special implementations of virtualized links.

Besides accounting for implementation characteristics, `adoptQoS` tracks paths and sub-paths, to consolidate requirements of multiple high-level links sharing the same resource-layer link. This happens when `link` is called multiple times for the same resource-layer link. Also, the information collected by `adoptQoS` is used to set up monitoring.

The result of the procedure are the application-layer links as resource-layer paths, where each path segment is a resource-layer link with a refined QoS attribute associated. The actual application of the attribute is resource allocation and deploying monitoring.

The implementation of links, resource allocation and QoS is technology specific. For example, in previous work [DgFKM11], we analysed that VMMs often use different metrics to model similar aspects of virtual and physical components. As QoS management heavily depends on the monitoring data obtained from the managed systems, QoS management must become ever more specific for the managed object. This can be seen as one of the main causes for the existing gap in QoS management, when introducing virtual components to IT infrastructures. The advantage of using the model and procedure introduced herein, lies in the generic structure and methodology for describing links and components. The model’s application results in a finite and sensible set of managed object types and corresponding management functions or strategies that allow for automated adoption, propagation and application of management information and functions from the high level view of virtual infrastructures to the concrete provisioning components.

## 4 Experiments

We apply our procedure to an exemplary VI with a single QoS attribute *data rate*. To illustrate the effect, we include a greedy “attacker” VM, which consumes resources to cause network performance degradation. With our procedure we identify the relevant components to perform QoS management to mitigate the attacker’s negative effect.

Figure 4 shows our test bed: a VM connected via a router to another VM, placed on our hosts `xensrv02`, `xensrv02` and `xensrv03`, respectively. While `xensrv01` and `xensrv02` are connected via a physical switch, `xensrv02` and `xensrv03` are directly linked. All VI components are volatile and only local switches are employed. The resource-layer topology is augmented by the classifications of its eleven path segments.

In this example, the QoS attribute is implemented correctly, when all path segments deliver at least the required data rate of 200 MBit/s ( $\approx 25$  MByte/s). The requirement is valid at

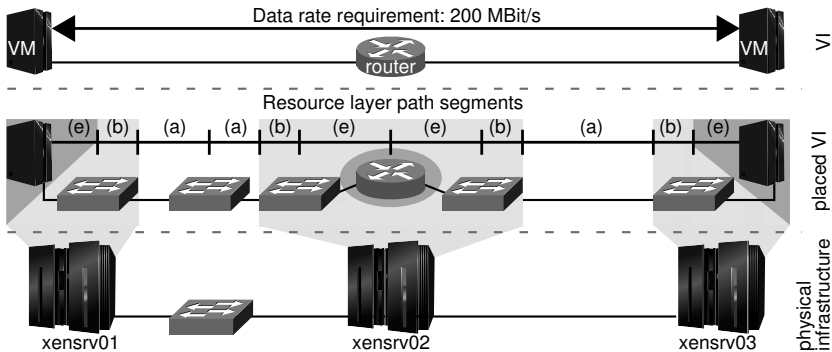


Figure 4: An exemplary VI mapped onto our test bed

application layer, and the data rate must be is adopted to correctly include the IP and Ethernet PCI. With standard Ethernet frames (up to 1500 bytes data), our use case specific `adoptQoS` evaluates  $\approx 200.3$  MBit/s for IP and  $\approx 200.6$  MBit/s Ethernet path segments.

When provided the two VMs as endpoints at layer 7, the procedure performs as follows:

1. Four iterations: OSI 17  $\rightarrow$  13. (a) Five iterations: OSI 17  $\rightarrow$  2
2. `adoptQoS`  $\rightarrow$  200.3 MBit/s (b) `adoptQoS`  $\rightarrow$  200.6 MBit/s
3. Twice: recursion into subpath (c) `link` for all found segments.

The application of our customized `adoptQoS` and `link` functions yields the following configuration steps:

- traffic shaping inside the VMs, for all traffic addressed to the other server,
- a new data path on each local switch for each direction and
- a new CPU-pool on each host for the VI's component.

We test our resulting configuration by comparing data flows between the servers with different kinds of load on `xensrv02`. To create flows, we use the program `mbuffer` to eliminate potential side effects due to disc I/O.

Our measurements are illustrated in Figure 5. Without the “attacker” VM, we achieve a fairly constant data rate of 39 MByte/s. Under CPU or network load, a drop below the required 25 MByte/s is observed. Having applied the developed configuration steps, the QoS goal is met, and the performance degradation evaded.

## 5 Related Work

The management challenge addressed in this paper pertains to network QoS in VEs. The special characteristics are that network components and network path segments are provisioned alongside endpoints one the one hand, and on-demand placement of virtual components which can be moved between hosts after initial provisioning on the other hand.

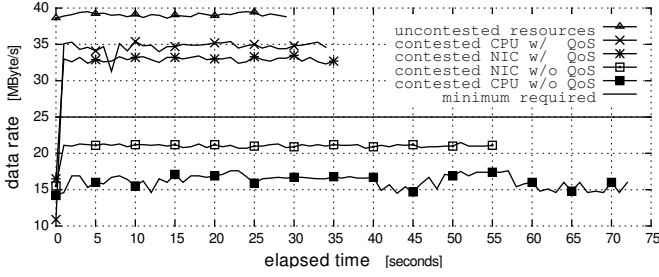


Figure 5: Achieved data rates

The gap between VI and component management has been addressed before. In [ea09c] and [ea09b] the authors suggest adaptive infrastructure management based on feedback loops. While this targets resources allocated to VMs, the network aspect is not handled specifically. The herein proposed approach can be used in such feedback loops and bridge the gap as described in this related work.

[ea09a] also identifies the need to adopt QoS attributes to technologies and (possibly hidden) network paths. This can be considered part of a solution, as they target networks only and the concept of volatile configurations does not fit their model.

Methods for refining resource allocations are still a topic for research, e.g. in [LN09]. This work provides the foundation for providing QoS links and paths in VEs. The recent study [ea13] analyses many such approaches and shows, that current VMMs don't "provide sufficient performance isolation [note: see [ea03]], to guarantee the effectiveness of resource sharing", which imposes a limit on efficiency and how accurate QoS can be enforced. The problem of performance of shared resources has been recognised before [Rix08] and may motivate an automated implementation of our procedure to address such problems and trigger dynamic adoptions of QoS attributes if combined with some of the available other approaches.

## 6 Conclusions

The proposed refinement procedure requires management information about the VI, i.e. application-layer topology, and its placement within the VE, i.e. resource-layer topology. It then reconstructs the placement of network links and in doing so filters the available management information for all required link segments and DTEs. The thereby generated information also contains all logical abstractions and compositions that contribute to the final "visible" link in the VI. This is the important information, required to adopt management information to sensible rules that can be applied.

The suggested method of generating views and analysing paths enables the detection of hidden paths and logical links and the triggering of special treatment, if required. With

this model as information base, generic rules for mapping VI QoS requirements to the resource-layer and the current state of a VE can be devised. This is a next step, out of the scope of this paper.

As part of an integrated management system for VEs, our approach is designed to allow substantial automation. It can be used to perform management on an existing VI, as demonstrated, or as part of the placement procedure. The specific `link` and `adoptQoS` functions can be built to provide feedback to acknowledge when VMs can be placed according to the specifications, or to force a different placement of VI components. Another aspect could be the actual planning in infrastructures where redundant network links allow for multiple network paths between hosts. To facilitate this heuristics can be added to the recursion stage of the procedure.

**Acknowledgments** The authors wish to thank the members of the Munich Network Management (MNM) Team for helpful discussions and valuable comments on previous versions of this paper.

## References

- [DgFKM11] V. Danciu, N. gentschen Felde, M. Kasch, and M. Metzker. Bottom-up harmonisation of management attributes describing hypervisors and virtual machines. In *Proceedings of the 5th International DMTF Workshop on Systems and Virtualization Management: Standards and the Cloud (SVM 2011)*, volume 2011, Paris, France, October 2011. Distributed Management Task Force (DMTF), IEEE Xplore.
- [ea98] Paul Ferguson et al. *Quality of service: delivering QoS on the Internet and in corporate networks*. Wiley New York, NY, 1998.
- [ea03] Paul Barham et al. Xen and the art of virtualization. *ACM SIGOPS Operating Systems Review*, 37(5):164–177, 2003.
- [ea09a] Karsten Oberle et al. Network virtualization: The missing piece. In *Intelligence in Next Generation Networks. ICIN 2009. 13th Int. Conf. on*, pages 1–6. IEEE, 2009.
- [ea09b] Pradeep Padala et al. Automated control of multiple virtualized resources. In *Proc. of the 4th ACM Europ. conf. on Computer systems*, pages 13–26. ACM, 2009.
- [ea09c] Qiang Li et al. Adaptive management of virtualized resources in cloud computing using feedback control. In *Information Science and Engineering (ICISE) 1st International Conference on*, pages 99–102. IEEE, 2009.
- [ea13] Yiduo Mei et al. Performance analysis of network I/O workloads in virtualized data centers. 2013.
- [JN04] Jingwen Jin and Klara Nahrstedt. QoS specification languages for distributed multimedia applications: A survey and taxonomy. *Multimedia, IEEE*, 11(3):74–87, 2004.
- [LN09] J Lakshmi and SK Nandy. I/O Device Virtualization in the multi-core era, a QoS perspective. In *Grid and Pervasive Computing Conference, 2009. GPC'09. Workshops at the*, pages 128–135. IEEE, 2009.
- [MD10] Martin G Metzker and Vitalian A Danciu. Towards end-to-end management of network QoS in virtualized infrastructures. In *Systems and Virtualization Management (SVM), 2010 4th Int. DMTF Academic Alliance Wksp. on*, pages 21–24. IEEE, 2010.
- [Rix08] Scot Rixner. Network virtualization: Breaking the performance barrier. *Queue*, 6(1):37, 2008.

# **IT-Sicherheit**



# Smart Defence: An Architecture for new Challenges to Cyber Security

Robert Koch, Mario Golling and Gabi Dreo Rodosek

Munich Network Management Team (MNM-Team)  
Universität der Bundeswehr München  
Werner-Heisenberg-Weg 39  
D-85577 Neubiberg  
{robert.koch, mario.golling, gabi.dreo}@unibw.de

**Abstract:** The last years have seen an unprecedented amount of attacks. Intrusions on IT-Systems are rising constantly - both from a quantitative as well as a qualitative point of view. Recent examples like the hack of the Sony Playstation Network or the compromise of RSA are just some examples of high-quality attack vectors. Since these Smart Attacks are specifically designed to permeate state of the art technologies, current systems like Intrusion Detection Systems (IDS) are failing to guarantee an adequate protection. In order to improve the protection, an analysis of these Smart Attacks in terms of underlying characteristics has to be performed to form a basis against those emerging threats.

Following these ideas, this paper starts by presenting individual facets of Smart Attacks in more detail. Inspired by the original definition of the term Advanced Persistent Threat of the Department of Defense, subsequently, the term Smart Attack is defined. Our architecture for Smart Defence focuses on three main elements: We propose the use of advanced geolocation for a geobased intrusion detection (e.g., inspecting new connections - originating from a location very close to where a recent attack was launched - more detailed than other connections). Furthermore, we will present our concepts on supervising Commercial Off-The Shelf (COTS) products (soft- and hardware), as both are nowadays used also in security environments. In addition, we will also show our concepts for similarity-based, multi-domain correlation as well as the corresponding proof-of-concept.

## 1 Introduction

The social and economic success of a society is increasingly dependent on a secure cyber space. Unfortunately, with the growing importance of the Internet for many areas within the last couple of years (e.g., for critical infrastructures such as energy or water supplies), securing the cyber space has not increased with the same speed. Quite the contrary, the number of attacks and their complexity is constantly increasing. The rising complexity, the variety of intelligence gathering techniques to access sensitive data, or the persistence and



stealthy characteristics of sophisticated recent attack vectors are just some of the reasons why state of the art security systems like (Next Generation-) Firewalls, Antivirus Systems or IDSs are failing to guarantee an adequate protection. Given these *Smart Attacks*, defense mechanisms need to be adapted to fill the gap. Therefore, it is the goal of this publication to develop new solutions for IT-security to face these new challenges. In this context, our architecture, specifically designed to defend against Smart Attacks, is presented within this publication.

The paper is structured as follows: Section 2 identifies characteristics of Smart Attacks and defines the term. The subsequent Section 3 examines related work. The architecture itself is described in Section 4. In Section 5 a prototype of the proposed detection approach is presented. Finally, Section 6 concludes the paper.

## 2 Characteristics of Smart Attacks

Over the past years, the sophistication of attacks has increased dramatically. This evolution of attacks is reflected by well-known terms like “Targeted Attacks” or “Advanced Persistent Threat” (APT). An attack can be considered as targeted if it is *intended for a specific person* or organization, typically *created to evade traditional security defenses* and frequently makes use of advanced social engineering techniques [Sym]. Wrt. the original definition given by the DoD, APTs are regularly originated by *nation state actors* or at least state-driven. However, the term APT is often misused for different kinds of highly professional attacks, which are more likely Targeted Attacks. APTs are designed to stay below the radar, and remain undetected for as long as possible; a characteristic that makes them especially effective, moving quietly and slowly in order to evade detection [Sym].

With the ongoing sophistication of attacks, these terms are not able to describe the special characteristics in a proper way. For example, high security networks often demand what is called an “air gap”, meaning, that Intranet and Internet have to be physically separated. However, already several examples are known in which the air gap has been overcome, since there is often still a necessity of a controlled data flow between the secured and the insecure network (“swivel chair interface”). The most prominent example of this is likely to be Stuxnet. After the recent emerging of some documents leaked by Snowden, describing technologies implemented and used by the NSA to bridge the air gap, e.g., “HOWLERMONKEY”, “SOMBERKNAVE” and “COTTONMOUTH”, this endangerment is increasing dramatically. In addition of being *targeted*, *sophisticated* and *persistent*, Smart Attacks are *camouflaged* (using alternative and special covert communication channels), *multilayered* (attacking firstly the perimeter networks if no direct attack is possible and afterwards the isolated networks). Another characteristic of Smart Attacks is that they are *interdisciplinary*: For example, specialists for Human Intelligence (HUMINT) are responsible for social engineering and information gathering, programmers develop a special kind of sophisticated attack code while service technicians install a new malicious

device in the network of a company (or as an “interdiction” process during the delivery).

Taking these new characteristics into account, *Smart Attacks* are defined as follows:

An attack is considered *smart* if it is aiming at a single target or a limited target group, which is exploited in-depth. The attack is executed via the combined, interdisciplinary exploitation of multiple domains in a camouflaged way, where the means and levels of exploitation of the individual domain is below the particular detection respectively suspicion thresholds. Smart Attacks can contain, or be limited to, sleeper exploitations which are (pre-)installed in software or hardware and waiting for a specific time or an external trigger for activation.

In particular, Smart Attacks are utilizing the exploitation of multiple domains without attracting attention in the single domain and can even be pre-installed: Therefore, a detection with current security systems is highly unlikely.

### 3 Related Work

To encounter the extensive threats opened up by today’s Smart Attacks, Smart Defence comprises numerous techniques: Effective classic protective measures as well as sophisticated new principles are required and have to be combined for an efficient protection. Because of the limited space of the publication, only a few examples of significant related work can be presented. Therefore, selected works of the most important areas, which are related with principles of our architecture, are discussed briefly as follows.

The basic principle of our architecture is the combination of established protection measures with our new developed components to achieve optimal detection results. In [EO10], the authors presented a collaborative intelligent intrusion detection system (CIIDS), which combines misuse- and anomaly-based systems. The resulting alerts of CIIDS are correlated and fuzzy logic is used to reduce the false alarm rate. While the idea of combining the advantages of the two detection principles sounds meaningful, no implementation or evaluation has been done by the authors. Even more, the DARPA 1999 IDS Evaluation dataset was proposed as being used for evaluation. While this dataset is well-known and has been used for numerous evaluations of IDSs, it also has been criticized widely because of the shortcomings wrt. data generation as well as the composition of the dataset (e.g., see [MC03]). An important possibility to identify outliers and therefore, abnormal and malign behavior, is the use of similarity measurements. Choi et al. give a comprehensive overview of binary similarity and distance measures [CCT10]. These methods can be used for intrusion detection. E.g., Yamada et al. proposed [Yam07] a system using similarity measures for the detection of malign access to servers within encrypted connections. In [FAH08], Foroushani et al. presented a system for intrusion detection in encrypted connections to servers secured

by the SSH protocol. While these systems are able to search for intrusions in encrypted network traffic, comprehensive configurations and server profiles are required, prohibiting an application at large.

A fundamental problem for the implementation of a reliable IDS is located very deep - the untrustworthiness of the hardware. E.g., see the debate in 2012 about hardware backdoors implemented in the Actel/Microsemi ProASIC3 FPGA or the discussion about the security of network products from ZTE and Huawei (e.g., see [RR12]). The disclosure of more and more details about the NSA wiretapping scandal underlines this problem area of sophisticated hardware backdoors in COTS products. An intensive discussion about the security of COTS began back in 2000 [NAT00], but only focuses on the security of software products. Some work has been done to improve the reliability of the IC fabrication and for the detection of hardware backdoors and trojan circuits (e.g., see [WP12]). However, the available techniques are still limited, only detecting special kinds of circuits or not being applicable in practice.

Another opportunity for improving detection results is the use of geo-information within the intrusion detection process. [KGR13] gives a comprehensive overview of the possibilities and restrictions of different approaches for the geolocation of IP addresses.

## **4 Smart Defence Architecture**

Because of their special characteristics, Smart Attacks can only be countered with Smart Defense. For this purpose, our architecture (which is depicted in Figure 1) makes use of the following three main principles: (1) perpetuation of successfully establish security concepts ("keep the tried and trusted") - gray color, (2) extension of local security mechanisms to include new elements (such as intrusion detection in encrypted environments, advanced geo-based intrusion detection, Layer 1 signal analysis) - orange color, and (3) use of external knowledge to improve the local detection results (from local to global knowledge) - blue color. Please note that a component for realizing the human intelligence knowledge, which is currently under development, is not depicted in the Figure.

### **Perpetuation of Well-Established Protective Measures**

Although current protection systems are not able to ensure an adequate protection against new challenges, it must not be forgotten that in the course of IT security research over the last 30 years, a high number of security measures has been developed that are now an integral component of any IT security strategy. Referring to the technical measures, among others, this includes [Nat13]: Switch-based options like fine-grained Access Control Lists (ACLs), port security and community/isolated port-based Virtual Local Area Networks (VLANs), firewalls, patch and vulnerability management servers (such as Windows Server Update Services), virus scanners and Network Intrusion Detection Systems (NIDS), e.g., the

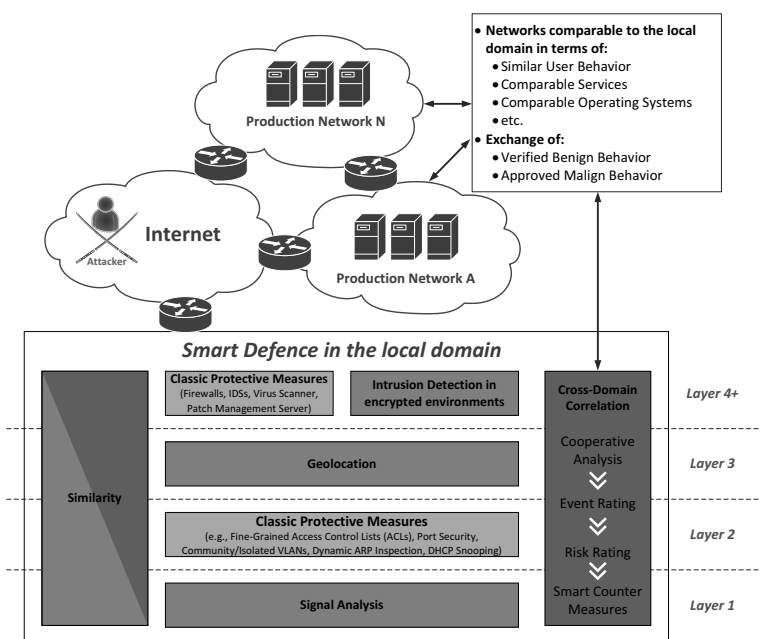


Figure 1: Overview of the architecture

signature-based system SNORT. While these techniques are not able to prevent sophisticated attacks, they can oppress most of the daily threats which account for approx. 80 percent of all attacks. Therefore, they build the base of the defence architecture.

### Integration of new Security Mechanisms

In order to address the new facets of Smart Attacks, well-established protective measures have to be enhanced with new security mechanisms (depicted in orange).

**Similarity:** Traditional forms of intrusion detection heavily rely on the recognition of pre-set patterns, which are an expression of previously acquired knowledge (usually about negative behavior). For example, IDSs like Snort or Antivirus Scanners are searching for known patterns (signatures). However, this involves in particular two disadvantages: (1) These systems are reactive by definition (i.e. the knowledge first needs to be collected) and (2) already slight variations can lead to the fact that a signature (representation of knowledge) is not applicable to a similar case. Even the use of heuristics cannot remedy this disadvantage sufficiently, because only simple variations (e.g., new members of a family of a known virus) can be detected or, with a broader detection spectrum, false alarm rates are rising significantly. In contrast to this, our architecture applies similarity-based detection principles on different layers and within different components. A variety of events and data can be analyzed based on similarity evaluations, e.g., the behavior of multiple users or the degree of likeness of two encrypted packet streams. The fundamental, anomaly-based

technique used for the detection is as follows: Similarity search based on a distance function represents a way of quantifying the proximity of two objects [ZADB06]. With regard to detection, similarity based detection generally focuses on either how certain traffic is similar to (known) anomalies or how different it is from normal behavior. Even more, because often the benign behavior prevails malign behavior by far, similarity measurements can be used to identify attacks without any knowledge about the benign behavior or the service under consideration. For example, most connections to a web server of a company will be benign, while a minor portion will be malign (e.g., brute-force login attempts or SQL injections). By calculating the similarity between the different connections, the malicious ones can be identified and isolated with high probabilities [Koc11].

*Signal Analysis:* COTS products have revolutionized IT. While COTS products are perfectly suited for the consumer market, this may be different in security critical environments, for instance due to the loss of control over the production process. Here, using COTS may result in introducing harmful new attack vectors into these systems. On the other hand, special developments are nowadays not financially feasible in most cases. Therefore, COTS products are also used in security critical areas. To improve the security of these devices and to enable a detection of manipulation as well as attacks, we introduce a new security concept in our architecture. Focusing on the hardware layer, the idea is to perform a Layer 1 signal analysis and compare the output of multiple systems built up of hardware components of different vendors (comparative systems). On Layer 1, the signal analysis is done by evaluating the frequency spectrum as well as measuring signal similarities and using multiple techniques like cross covariance and cross correlation functions. By that, distances (*similarity*) of different flow and packet characteristics are evaluated for the detection of manipulated COTS hardware components. E.g., differing behavior or additional traffic (which will be typically highly suspicious if only one or a limited number of comparative systems exhibit this behavior) of specific hardware components can be detected and further investigations can be initiated. For a more comprehensive overview, take a look at [KGR14].

*Geolocation:* Similarity with regard to the origination of IP-addresses (*geolocation*) can be applied for intrusion detection as well. E.g., a recent study from the security company Mandiant - claiming to analyze Chinas Cyber Espionage Units - proclaimed, “a large share of hacking activity targeting the US could be traced to an office building in Shanghai”. Here, similar to greylisting in emails, geolocation allows to (i) correlate attacks detected with new connections and (ii) as a consequence to classify traffic a priori as more suspicious (thus particularly allowing to inspect this traffic in more detail). While suspicious IP addresses (e.g., proxies, TOR exit nodes or IPs marked by blacklists) can be blocked quite easily, this is not enough. For example, a hacker can use systems respectively IP addresses of an infiltrated network (stepping stone) to execute attacks onto the domain. Therefore, if an address is identified as suspicious, the IP as well as the addresses can be temporarily blocked to hamper the possibility of the attacker to use different IPs of the infiltrated network for the attack.

*Similarity-Based Intrusion Detection:* Today's IDSs either need large databases (signatures/

patterns) in case of signature and payload-based detection methods or often a learning phase to analyze the normal behavior of a network. In contrast to these techniques, our approach makes use of inherent knowledge of systems and services. For example, normally the overall majority of all connections to a website (e.g., a Web shop) will be benign, while only a little fraction will be malicious (attacking the Web shop, etc.). Due to the fact that the different actions - benign and malicious - are quite different concerning specific characteristics (like number of network packets, delay between the individual packets or number of individual connections), the majority of connections and their similarity can be used to separate the benign traffic from the malicious even without having to know - in advance - how the benign behavior looks like. Following the idea of a majority decision, a separation can be performed in real-time, without complex configurations and without the need to know specific details about the system/environment.

### **Cross-Domain Correlation**

In particular, the last idea can also be extended to other domains. If the networks of the foreign domain have a high similarity (in terms of network and user behavior, services provided, operating systems installed, etc.), this knowledge can be used to extend the pool on which the comparative analysis is performed. In addition, the alerts of the different IDSs can be exchanged with each other in real time. Here, an analysis is first performed locally (*Event Rating*), to differentiate whether an alarm is present and if so, the corresponding severity. However, on the side of the individual IDS, diverse forms of rating models exist. In order to be applicable to a wide range of systems, the first step comprises the transition of the proprietary Risk Rating Models into the Common Vulnerability Scoring System (CVSS) by applying a transition formula. The advantage is that now the individual alerts are comparable and so a comparative analysis of the risk assessment is made. This in turn is a prerequisite for smart countermeasures (as for instance an earlier notification of the operator or automated reactions).

## **5 Prototype**

To verify the detection capabilities of our architecture, a proof-of-concept is currently under development. Not all modules are completely implemented yet, but most parts are ready and under extensive evaluation. At the moment, the modules are running in stand-alone mode; isolated evaluations are presented as follows. The complete architecture will focus on the detection of sophisticated Smart Attacks. Due to the usage of real-time measurements to identify malign behavior, an adaption of attackers to the detection scheme will be very challenging as long as the normal and the benign behavior has some distinctions (which should normally be the case). After the completion and integration of the final components, a comprehensive overall evaluation including detection capabilities and limitations will be done and real-time aspects and reactions will be discussed in-depth. Because of the limited space of the publication, only some results are presented as follows, namely for

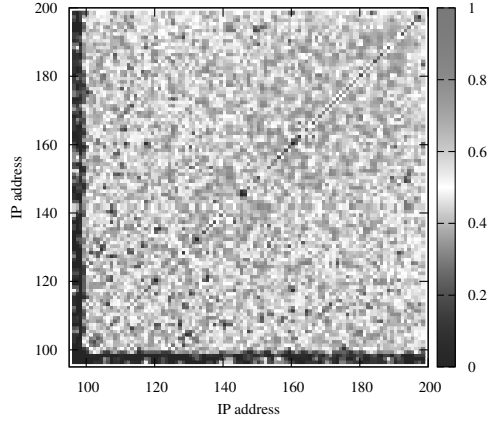


Figure 2: Similarity calculation of user behavior visualized by a colored map

the geolocation and the similarity module. The use of similarity is a basic functional principle throughout our architecture to, e.g., identify malicious behavior or the evaluation of user behavior. Based on the similarity measurements within (i) the local domain but also by (ii) the inclusion of knowledge from other domains, incidents can be checked and malicious behavior can be identified. See Fig. 2 for an example. The presented colored map was generated when encrypted connections of different users, accessing a server in the production network, were evaluated. While 99% of the users had been benign, 1% of the connections had been attacks executed against the server (brute force and SQL injection attacks). The streams of network packets of different users are correlated with each other to calculate the similarity of the data streams. While a correlation value of zero represents no similarity, increasing values show higher similarities. For a better illustration, IP addresses of attackers are below 100 (last octet), while benign users have addresses from 100 up to 200. Since this assignment of IP addresses is not used within our architecture, it does not influence the results, but simplifies the presentation. As shown in Fig. 2, there is a strong blue stripe on the left and at the bottom, representing low correlation values of the malicious connections. On the other hand, benign connections have higher correlation values, illustrated by red color: Because of the differences within the statistical features of benign respectively malicious data streams, benign sessions correlated with each other result in higher similarity values, while malicious sessions correlated with benign ones result in low values. Having 1% of malign connections, the Probability of Detection (PD) is 84.70% for a *single* evaluation. For the actual classification of a connection, multiple evaluations are used to increase the accurateness: Every connection is correlated with 11 other connections as this value was identified as the optimal number of correlation partners by empirical test runs. Based on that, the connections are classified: Having 1% of malicious connections, the final detection rate is about 99.3%. Therefore, our technique can be used to identify malign user behavior even without any knowledge about the used service, the transferred data (no need for decryption), a behavior model of the network or signatures. Note, that the system also doesn't need a configuration, in contrast to current systems using similarity for intrusion detection in encrypted communication (see Section 3).

Next, some results of the evaluation of the module for geolocation are given. Tab. 1 gives the detection accuracy of our IPv4 geolocation module for country and for city level. On country level, 99.78% of all addresses are identified correctly while on city level, our module is able to locate 90.49% of the addresses correctly. This outperforms the current other approaches, where the best one has an accuracy of 98% on country and 87.4% on city level (see [KGR13]).

Table 1: Accuracy of the geolocation module

| LOCALISATION LEVEL | ACCURACY |
|--------------------|----------|
| Country            | 99,78%   |
| City               | 90,49%   |

With an increasing usage of IPv6, the significance of an adequate IP geolocation will be even more important for the area of intrusion detection: While with IPv4, numerous private networks are used and masked by gateway and proxy IP addresses, the concept of IPv6 enables the possibility of identifying every device on the Internet because of the scheme for the address generation, using the hardware address of a device for building the IPv6 address. Even if Privacy Extensions (PE) are used for a randomized generation of the IPv6 address, at least the network part remains and can be used for identification. On the one hand, the IPv6 addressing scheme can be used to improve the geolocation-based portion of the proposed IDS architecture. On the other hand, because of the enormous number of possible IPv6 addresses, the actual geolocation process becomes much more challenging than in the case of IPv4. Therefore, the integration of an adequate IPv6 geolocation into our architecture will be part of our future work.

## 6 Conclusions and Outlook

Cyber attacks are increasing in complexity, persistence and stealthiness. APTs have more and more evolved towards so-called Smart Attacks. The paper defines the term Smart Attack and discusses characteristics, which are typically for this new very sophisticated kind of endangerment. Based on this, we propose a new architecture for Smart Defence, designed to cope with the special characteristics of this new generation of attacks. Geolocation and similarity-based intrusion detection are central components within our architecture. Based on similarity calculations, malign behavior (and therefore connections and users) can be identified without signatures or knowledge about the monitored services or characteristics of the traffic. The promising results of the evaluation of the prototypical implementations of the different modules on real data undermine the conceptual work. At the moment, the central management, control and coordination module has been implemented as well as the geolocation module. Our next steps will be to include the exchange of external domains in a better way as well as to concentrate more on Layer 1 signal analysis.



## Acknowledgements

This work was partly funded by Flamingo, a Network of Excellence project (ICT-318488) supported by the European Commission under its Seventh Framework Programme.

## References

- [CCT10] Seung-Seok Choi, Sung-Hyuk Cha, and C Tappert. A survey of binary similarity and distance measures. *Journal of Systemics, Cybernetics and Informatics*, 8(1):43–48, 2010.
- [EO10] H.T. Elshoush and I.M. Osman. Reducing false positives through fuzzy alert correlation in collaborative intelligent intrusion detection systems; A review. In *Fuzzy Systems (FUZZ), 2010 IEEE International Conference on*, pages 1–8, 2010.
- [FAH08] Vahid Aghaei Ferooshani, Fazlollah Adibnia, and Elham Hojati. Intrusion detection in encrypted accesses with SSH protocol to network public servers. 2008.
- [KGR13] Robert Koch, Mario Golling, and Gabi Dreo Rodosek. Advanced Geolocation of IP Addresses. In *International Conference on Communication and Network Security (ICCNS)*, pages 1–10, 2013.
- [KGR14] Robert Koch, Mario Golling, and Gabi Dreo Rodosek. Using Layer 1 Signal Analysis for the Supervision of COTS Products. In *9th International Conference on Cyber Warfare and Security ICCWS-2014*, 2014.
- [Koc11] Robert Koch. *Systemarchitektur zur Ein-und Ausbruchserkennung in verschlüsselten Umgebungen*. BoD-Books on Demand, 2011.
- [MC03] Matthew V Mahoney and Philip K Chan. An analysis of the 1999 DARPA/Lincoln Laboratory evaluation data for network anomaly detection. In *Recent Advances in Intrusion Detection*, pages 220–237. Springer, 2003.
- [NAT00] NATO Research and Technology Organisation. Commercial Off-the-Shelf Products in Defence Applications (The Ruthless Pursuit of COTS). NATO, 2000.
- [Nat13] National Institute of Standards and Technology. Security and Privacy Controls for Federal Information Systems and Organizations. *NIST Special Publication 800-53, Rev 4*, 2013.
- [RR12] Mike Rogers and Dutch Rupperberger. Investigative Report on the U.S. National Security Issues Posed by Chinese Telecommunications Companies Huawei and ZTE. Technical report, U.S. HOR, 2012.
- [Sym] Symantec Corporation. Industrial Espionage: Targeted Attacks and Advanced Persistent Threats (APTs).
- [WP12] Sheng Wei and Miodrag Potkonjak. Scalable hardware Trojan diagnosis. *Very Large Scale Integration (VLSI) Systems, IEEE Transactions on*, 20(6):1049–1057, 2012.
- [Yam07] Akira et al. Yamada. Intrusion detection for encrypted web accesses. In *Advanced Information Networking and Applications Workshops, 2007, AINAW'07. 21st International Conference on*. IEEE, 2007.
- [ZADB06] Pavel Zezula, G Amato, V Dohnal, and M Batko. *Similarity search*. Springer, 2006.

# Herausforderungen und Anforderungen für ein organisationsweites Notfall-Passwortmanagement

Felix von Eye, Wolfgang Hommel, Helmut Reiser

Munich Network Management Team  
Leibniz-Rechenzentrum (LRZ)  
Boltzmannstr. 1  
85748 Garching b. München  
{voneye,hommel,reiser}@lrz.de

**Abstract:** Das Hinterlegen von Passwörtern ist ein notwendiges Übel, wenn in Ausnahmesituationen auch solche Personen auf Systeme und Dienste zugreifen können sollen, die im regulären Betrieb nicht damit befasst sind. Bei einer großen Anzahl verwalteter Passwörter, die zudem häufig geändert werden müssen, skalieren herkömmliche Ansätze wie versiegelte Passwortzettel im Tresor nicht ausreichend. Dedizierte Softwarelösungen bieten Chancen, bergen aber auch Sicherheitsrisiken. Am Beispiel des Leibniz-Rechenzentrums stellen wir Herausforderungen, Anforderungen, Lösungsansätze, bisherige Betriebserfahrungen und offene Wünsche vor.

## 1 Einleitung

Die Eingabe einer Kombination aus Loginname und zugehörigem Passwort stellt in IT-Umgebungen mit normalem Schutzbedarf seit langer Zeit die am häufigsten eingesetzte Form der Authentisierung dar. Trotz Alternativen wie Tokens, Smartcards, elektronischen Ausweisen und verschiedensten Varianten der Biometrie, deren Vor- und Nachteile gegenüber der Passwort-Authentifizierung schon oft und lange diskutiert wurden (z. B. in [RZ06] oder [Mau08]), führt im Hochschulumfeld fast kein Weg an Passwörtern vorbei: Relativ geringe Implementierungskosten, ausreichend hohe Nutzerakzeptanz und die Seltenheit signifikanter Sicherheitsvorfälle legen nahe, dass Passwörter auch in Zukunft ein fester Bestandteil des täglichen Umgangs mit IT-Diensten sein werden.

Etwas vereinfacht dargestellt muss bei Passwörtern unterschieden werden, ob sie

- *persönliche* Zugangsdaten darstellen, beispielsweise zum Login an PCs, zum Zugriff auf die eigene Mailbox oder zur Anmeldung bei einem webbasierten Dienst, oder
- *mehreren Personen* bekannt sein sollen, z. B. root- bzw. Administrator-Passwörter für Server bzw. für das Management von IT-Diensten.

Für den ersten dieser beiden Fälle wird aus guten Gründen empfohlen, für unterschiedliche

Dienste bzw. Dienstleister verschiedene Passwörter zu verwenden (siehe z. B. [Bun11]). Die recht große Anzahl insbesondere an Webdiensten – durchschnittlich verwaltet jeder Nutzer im beruflichen wie auch im privaten Umfeld 25 verschiedene Accounts [FH07] –, motiviert den Einsatz von Passwortmanagement-Werkzeugen, die sich vom Passwortzettel in der Brieftasche über das Speichern im Browser bis hin zum Einsatz dedizierter Software erstrecken, wobei die Sicherheit dieser Lösungen regelmäßig geprüft werden sollte [BS12].

Auch für die System- und Dienstadministration empfiehlt sich generell die Verwendung persönlicher Zugangsdaten, also beispielsweise der Einsatz mehrerer personalisierter Kennungen mit root- bzw. Administratorenberechtigung auf Servern, wenn mehrköpfige Betriebsgruppen Zugriff haben sollen – alleine schon, um im Zweifelsfall einfacher nachvollziehen zu können, wer eine bestimmte Änderung an einem System vorgenommen hat. Dieser Teilaspekt der Benutzer- und Berechtigungsverwaltung wird gemeinhin als *Privileged Account Management* bezeichnet.

Das *Notfall-Passwortmanagement* deckt hingegen Situationen ab, in denen am Dienstbetrieb nicht direkt beteiligte Dritte in Ausnahmesituationen Zugriff auf Systeme erhalten sollen, um größeren Schaden abzuwenden. Beispielsweise kann es ein Ziel sein, dass ein rund um die Uhr besetztes Security-Team Zugriff auf möglicherweise kompromittierte Server erhält, auch wenn der zuständige Systemadministrator und z. B. seine Urlaubsvertretung längere Zeit nicht erreichbar sind. Zu den klassischen Lösungsansätzen für dieses Problemfeld gehören z. B. auf Zetteln notierte Passwörter in versiegelten Kuverts, die in einem Tresor deponiert und beim Eintreten eines Notfalls einem dazu Berechtigten ausgehändigt werden.

Solche Lösungen, die auf einer physischen Absicherung basieren, skalieren allerdings nur unzureichend: Anders als früher müssen nicht mehr nur einige Handvoll, sondern Dutzende von Passwörtern geeignet hinterlegt werden. Passwortrichtlinien sehen nicht nur regelmäßige Passwortänderungen vor, sondern erfordern diese auch beispielsweise bei Personalfuktuation – diese ist aufgrund der vielen befristeten Stellen im Hochschulumfeld inzwischen eher die Regel als die Ausnahme. Das Pflegen hinterlegter Passwortzettel entwickelt sich dadurch von der subjektiv lästigen Pflichtübung hin zum objektiven Zeitfresser.

Die Umstellung des Notfall-Passwortmanagements auf eine zentrale, serverbasierte Software-Lösung, die von jedem Administrator-Arbeitsplatz aus genutzt werden kann, birgt viele Chancen, aber auch viele (Sicherheits-) Risiken und will aufgrund ihrer Tragweite gut überlegt sein. In diesem Artikel stellen wir die Ergebnisse eines rund einjährigen Projekts am Leibniz-Rechenzentrum (LRZ) vor, in dem zunächst eine umfassende Bedarfs- und Anforderungsanalyse durchgeführt wurde, die in Abschnitt 2 zusammengefasst ist und zum Teil auf Vorarbeiten einer an der Ludwig-Maximilians-Universität München erstellten Bachelorarbeit [Gem12] basiert und auf deren Basis mehrere kommerzielle Softwarelösungen evaluiert wurden. In Abschnitt 3 beschreiben wir die implementierte Lösung, unsere Vorgehensweise bei der Einführung und Migration sowie die bisherigen Praxiserfahrungen. Abschließend fassen wir in Abschnitt 4 das bislang Erreichte zusammen, skizzieren noch offene Punkte für die Weiterentwicklung und beurteilen die Übertragbarkeit auf andere Hochschulrechenzentren.

## 2 Entwicklung eines Kriterienkatalogs

Die zentrale Anforderung an eine Passwortmanagement-Lösung ist, dass Passwörter, die der Lösung anvertraut werden, sicher und jederzeit abrufbar abgesichert werden, wobei sicher bedeutet, dass zu jedem Zeitpunkt sichergestellt sein muss, dass Unberechtigte nicht in der Lage sind, mit vertretbarem Aufwand die Passwörter im Klartext zu erhalten. Beim Notfall-Passwortmanagement wird die zentrale Anforderung des Passwortmanagements dahingegen erweitert, dass der Kreis der zum Einsehen eines Passwordeintrags berechtigten Personen temporär vergrößert wird. Dabei besteht jedoch die Herausforderung, dass ein potentieller Missbrauch einer Notfall-Berechtigung weitestgehend ausgeschlossen wird oder wenigstens eine zeitnahe Alarmierung stattfindet.

Die Kriterien, nach denen man eine Notfall-Passwortmanagement-Lösung bewerten kann, kann man grob in die Kategorien *Architektur*, *Passworterstellung*, *Notfallautorisierung*, *Auditierung* und *Sicherheitsanforderungen* einteilen. Im Folgenden werden diese Kategorien beschrieben und die darin enthaltenen Kriterien erläutert.

### 2.1 Architektur

Ein organisationsweites Notfall-Passwortmanagement sollte, wenn organisatorisch möglich, so viele der insgesamt eingesetzten und nicht nur persönlich genutzten Passwörter wie möglich enthalten. Dies dient dazu, dass in einer Notfallsituation keine unnötigen Verzögerungen entstehen, da zu jedem Passwordeintrag bekannt sein sollte, wo er zu finden ist. In vielen Fällen bietet es sich daher an, die Passwörter in einer zentralen Passwortmanagement-Instanz abzulegen und die Autorisierung innerhalb dieser vorzunehmen; es ist jedoch ebenso möglich, dass mehrere Instanzen parallel oder hierarchisch betrieben werden.

Vor Einführung eines Notfall-Passwortmanagements ist somit zunächst zu ermitteln, welche zentrale Struktur man für das Speichern von Passwörtern vorsieht. Erst wenn es eine einheitliche Struktur gibt, können Passwordeinträge von Personen aufgefunden werden, die diese Passwordeinträge nicht angelegt haben, vor allem, wenn sich die Passwörter über mehrere Datenbanken oder Instanzen erstrecken.

Damit man eine wirksame Autorisierung und Authentifizierung durchführen kann, ist es notwendig, dass eine organisationsweite Passwortmanagement-Lösung mehrbenutzerfähig ist. Die populären und meist für den privaten Einsatz konzipierten Passwortmanager wie KeePass, KWallet oder die in modernen Browsern integrierten Passwortmanager sind üblicherweise nicht in der Lage, mehrere voneinander unterschiedliche Nutzer mit zwar unterschiedlichen Berechtigungen, aber einem gemeinsamen Datenbestand zu unterstützen. Da das Anlegen und Verwalten von Benutzern jedoch kein Teil des Passwortsondern des Identity-Managements ist, ist eine Integration mit Verzeichnisdiensten wie LDAP oder Active Directory erforderlich.

Ebenfalls für den Aufbau des Passwortmanagement-systems ist von Interesse, wie die Absicherung gegen diverse mögliche Ausfälle oder Fehler realisiert ist. Da ein Notfall-Pass-

wortmanagement als Spezialfall des normalen Passwortmanagements aber genau für diese Fälle eingesetzt werden soll und man somit nicht davon ausgehen kann, dass die hauptamtlichen Administratoren zur Verfügung stehen – beispielsweise könnte dies während einer Rufbereitschaft außerhalb der üblichen Bürozeiten sein – ist die Zugänglichkeit des Systems von entscheidender Bedeutung. Üblicherweise ist jedoch die Ausfallsicherung weniger eine Problemstellung des technischen als eher des organisatorischen Systems. Wenn beispielsweise das Notfall-Passwortmanagement von der Existenz und Funktionsfähigkeit anderer Systeme abhängt, so wird es nach dem Ausfall eben dieser Systeme nicht mehr funktionieren. Unabhängig davon muss jedoch eine gewisse Redundanz unterstützt werden.

Ferner sollte ein Notfall-Passwortmanagement-System so in der Infrastruktur platziert werden, dass dieses über klar definierte Zugriffswege erreichbar ist. Weiterhin hat jegliche Kommunikation und jeglicher Eintrag innerhalb des Systems verschlüsselt zu erfolgen, so dass selbst im Falle der (ggf. auch physischen) Kompromittierung des Systems die gespeicherten Passwörter nicht ohne zusätzlichen Aufwand in falsche Hände gelangen können.

Was ein Notfall-Passwortmanagement-System noch abrundet, aber nur in wenigen auf dem Markt erhältlichen Produkten zu finden ist, ist die Möglichkeit, dass das Notfall-Passwortmanagement-System Passwörter automatisch auf Zielsystemen ändert und diese dann in der eigenen Datenbank aktualisiert. Zwar ist diese Funktion nicht unproblematisch, da durch eine falsch konfigurierte Befehlsfolge schnell aus einer harmlosen Passwortänderung eine Systemstörung wird, aber es eröffnet die Möglichkeit, dass ein Missbrauch der Notfall-Berechtigung durch eine zeitliche Einschränkung verringert wird. Diese automatisierte Passwortänderung wäre somit einer Forderung nach einer regelmäßigen Änderung von Passwörtern, wie dies beispielsweise der BSI Grundsatz [Bun] vorsieht, gleichzusetzen.

## **2.2 Erstellen der Passwörter**

Die Nutzung einer Passwortmanagement-Lösung bringt für Administratoren den Vorteil, dass nicht mehr jedes einzelne Passwort für jeden einzelnen Dienst individuell gewählt, anhand einer möglicherweise vorhandenen Passwortrichtlinie verifiziert und manuell an eine Passwortänderung gedacht werden muss. Diese Arbeitserleichterung setzt jedoch voraus, dass Passwortgeneratoren zur Verfügung stehen, die beliebige Passwortrichtlinien unterstützen und ausreichend zufällige Passwörter erstellen können.

Eine beliebte Möglichkeit, sich einer regelmäßigen Änderung von Passwörtern zu widersetzen, von Administratoren ist, zwei Passwörter zu haben, die immer im Wechsel verwendet werden. Sind diese Passwörter genügend unterschiedlich, so überlisten sie ein einfaches Prüfsystem, das im Allgemeinen nur das aktuelle mit dem neuen Passwort vergleicht und bewertet, ob das neue Passwort genügend Unterschiede aufweist, um als „neues“ Passwort zu gelten. Somit sollte eine Passwortmanagement-Lösung ein Log der bisher verwendeten Passwörter enthalten, die bei solch einer Prüfung herangezogen werden kann. Diese History der bereits verwendeten Passwörter muss jedoch genauso gut abgesichert sein wie

das aktuelle Passwort, um mögliche Sicherheitsrisiken auszuschließen.

Auch an anderer Stelle ist eine solche History der Passworteinträge interessant. Im Fall eines Restore eines älteren Backups kann es vorkommen, dass die regelmäßige Passwortänderung des Systems bereits auf Grund einer oder mehrerer Passwortänderungen ein neues Passwort gespeichert hat. Ist eine solche History nicht vorhanden, würde sich der Zugriff auf das Backup als schwierig erweisen bzw. es wäre nötig, bei jedem Backup darauf zu achten, dass wichtige Systempasswörter in einzelnen Passworteinträgen von Hand ergänzt oder extern gespeichert werden.

Weiterhin muss es einem Nutzer des Passwortmanagements möglich sein, zu bestimmen, wer seinen Passworteintrag einsehen darf und wer nicht. Vor allem beim Notfall-Passwortmanagement muss diese Berechtigungsverwaltung noch feingranularer sein als dies beim normalen Passwortmanagement möglich ist. So ist eine einfache „Lesen“-Freigabe für ein Notfall-Passwortmanagement nicht ausreichend, da unterschieden werden muss, ob ein lesender Zugriff im Rahmen der täglichen Arbeit oder ob dieser lesende Zugriff im Rahmen einer temporären Notfall-Berechtigung stattfindet. Üblicherweise sollen Personen, die über eine temporäre Notfall-Berechtigung verfügen, nicht in der Lage sein, die Einträge außerhalb fest vorgegebener Umstände einzusehen, deren Eintreten zu dokumentieren ist.

Neben dem Einsehen eines Passworteintrages sollte ein Nutzer auch in der Lage sein, zu bestimmen, welche Einträge im Ganzen für wen sichtbar sind und welche anderen Benutzern verborgen bleiben. Gerade wenn Einträge, für die man überhaupt keine Berechtigung hat, ausgeblendet werden, wird die Übersichtlichkeit und Nutzbarkeit des Systems erhöht.

## **2.3 Autorisierung einer temporären Zugriffsberechtigung**

Zentraler Punkt eines Notfall-Passwortmanagements ist das Bereitstellen von temporären Zugriffsberechtigungen für den Fall, dass ein Zugriff auf ein System von jemanden durchgeführt werden muss, der ansonsten nicht zum Kreis der Administratoren des Systems gehört. Da diesem im Allgemeinen das Passwort für den Zugriff nicht bekannt sein sollte, ist für ein erfolgreiches Anmelden eine aktuelle Zugriffsberechtigung nötig.

Wie ein solches Ausweiten der normalen Berechtigung technisch und organisatorisch umgesetzt werden sollte, ist von Organisation zu Organisation unterschiedlich. Denkbar wären beispielsweise versiegelte Passworteinträge, die nicht unbemerkt zu öffnen sind. Erst wenn der Fall eintritt, dass eine temporäre Zugriffsberechtigung benötigt wird, wird das Siegel gebrochen und der Zugriff auf das Passwort ist möglich. Selbstverständlich muss in diesem Fall eine Benachrichtigung an den jeweils zuständigen Administrator gesandt werden, wobei sich dafür beispielsweise E-Mails, SMS oder automatisierte Sprachanrufe eignen.

Da sich diese Art der Privilegieneskalation weder für alle Einsatzszenarien noch für alle Organisationsformen eignet, da es sich hierbei um eine Tätigkeit handelt, die viel mit Vertrauen und wenig mit Kontrolle zu tun hat, bietet es sich an, auch andere Formen der temporären Zugriffsberechtigung zu fordern wie beispielsweise Mehr-Augen-Prinzip oder explizite Berechtigungspfade, das heißt, dass es explizit genannte Personen gibt, die nacheinander ihr Einverständnis zur Privilegieneskalation geben müssen. Diese Methoden eig-

nen sich, um einen potentiellen Missbrauch der temporären Zugriffsberechtigung wirksam zu unterbinden, da ein Mitarbeiter alleine dazu nicht die Berechtigung besitzt. Ist jedoch in einer Notfallsituation nur eine der nötigen Personen, die ihr Einverständnis geben müssen, nicht verfügbar, so kann in einer solchen Situation auf den Passworteintrag nicht zugegriffen werden, so dass organisatorisch immer sichergestellt werden muss, dass jederzeit ausreichend viele verantwortlichen Entscheider verfügbar sind.

Ein wichtiger Punkt bei einer temporären Zugriffsberechtigung ist, dass sie temporär ist, das heißt, dass eine einmalig gewährte Zugriffsberechtigung nicht dauerhaft gültig sein sollte, sondern nach einer fest oder dynamisch vorgegebenen Zeit wieder entzogen wird. Dabei ist es wichtig, dass dieser Vorgang durch die Notfall-Passwortmanagement-Lösung durch einen Automatismus oder wenigstens durch einen Teilautomatismus unterstützt wird, da ansonsten im täglichen Betrieb untergeht, dass beispielsweise ein Passworteintrag entsiegelt wurde. Ebenfalls denkbar, aber technisch viel aufwändiger ist die schon oben erwähnte automatische Passwortänderung nach dem Gewähren einer temporären Zugriffsberechtigung. Dies verhindert, dass man das Wissen um ein Passwort auch nach dem Entzug der Zugriffsberechtigung weiter nutzen kann.

## **2.4 Audit**

Gerade im Hinblick auf die temporäre Zugriffsberechtigung ist es nötig, je nach Schutzbedarf der Einträge mitzuprotokollieren, wer welche Passworteinträge einsieht bzw. welche Aktionen mit einem Passworteintrag durchgeführt wurden. Dabei sollte, soweit dies möglich ist, das Log auditsicher, d.h. resistent gegenüber Manipulationen sein, so dass im Idealfall nicht einmal Tool-Administratoren der Lösung die Möglichkeit haben, alte Logeinträge zu manipulieren.

In vielen Einsatzszenarien ist es weiterhin hilfreich, wenn Logeinträge nicht nur innerhalb der Notfall-Passwortmanagement-Lösung zu finden sind, sondern wenn es einen automatischen Versand von Benachrichtigungen gibt, der die zuständigen Administratoren über wichtige Ereignisse informiert. Neben der temporären Zugriffsberechtigung kann es allgemein für das Monitoring eines Systems interessant sein, zu ermitteln, wer sich zu welchem Zeitpunkt ein Passworteintrag ansieht, um im Anschluss auf das System zuzugreifen.

## **2.5 Sicherheitsanforderungen**

Gerade bei der zentralen Sammlung und Verwaltung von Passwörtern für administrative Zugänge auf Systeme muss man sensibel auf die Sicherheitsbedenken der Administratoren eingehen, da diese sich sehr stark auf die Bereitschaft der Administratoren auswirken, dem Notfall-Passwortmanagement-System Zugangsdaten anzuvertrauen. Aber neben den subjektiven Sicherheitsbedenken einzelner Administratoren muss man sich jedoch vor Einführung eines Notfall-Passwortmanagements durchaus bewusst sein, dass ein zentraler Ort, an dem ein Großteil der Passwörter und sonstigen Zugriffsberechtigungen einer Or-

ganisation abgelegt ist, bei Bekanntwerden tatsächlich ein attraktives Ziel für potentielle Angreifer darstellt. Aus diesem Grund ist es wichtig, im Vorfeld zu untersuchen, ob die gewählte Softwarelösung und das System, auf der diese betrieben wird, den eigenen Sicherheitsanforderungen genügt oder ob es möglich bzw. nötig ist, zusätzliches dediziertes Sicherheitsmonitoring zu installieren.

Viele Hersteller versprechen in ihren Produktbeschreibungen, dass jegliche Einträge mit den besten und sichersten kryptographischen Algorithmen verschlüsselt in der Datenbank abgespeichert werden. Der Nachweis dieser Behauptung ist jedoch selbst in einer Evaluation von Software-Lösungen nur sehr schwer bis gar nicht zu erbringen und auch bei Open-Source-Lösungen nur mit unverhältnismäßigem Aufwand nachweisbar, dass keine Bugs oder Backdoors vorhanden sind. Zudem sollte berücksichtigt werden, dass die verschlüsselte Speicherung der Passworteinträge nur einer von mehreren wichtigen Sicherheitsbausteinen ist. Mindestens genauso wichtig sind die zuverlässige Implementierung des Berechtigungsmanagements und dem Schutzbedarf angemessene Authentifizierungsmöglichkeiten – beispielsweise beugt eine Zwei-Faktor-Authentifizierung dem Problem vor, dass ein vom Angreifer ausgespähtes Benutzerpasswort den Zugriff auf alle seine Passworteinträge ermöglicht.

### **3 Inbetriebnahme eines Notfall-Passwortmanagements**

Bei der Einführung einer neuen Software-Lösung beginnt man im Allgemeinen nicht auf der grünen Wiese, sondern hat schon eine vorhandene Infrastruktur, in die sich die neue Software-Lösung integrieren muss. Auch am LRZ wurden in der Vergangenheit Notfall-Passwörter in verschiedenen Arten und Weisen aufbewahrt, wobei sich die Aufbewahrung von Fall zu Fall unterschieden hat, jedoch waren die gewählten Lösungen zueinander inkompatibel und zumeist offline. Aus diesem Grund wurde vor etwa eineinhalb Jahren beschlossen, dass ein hausweites Notfall-Passwortmanagement eingeführt werden soll.

Zu diesem Zweck wurden einige Tools und Konzepte evaluiert, die anhand des oben genannten Kriterienkatalogs bewertet wurden. Dabei hat sich herausgestellt, dass alle (Personal) Passwortmanager für ein Notfall-Passwortmanagement ungeeignet sind, da die dafür nötige Freigabe von Berechtigungen nicht zeitnah und ohne weiteres umgesetzt werden kann. Da dies am LRZ die grundlegende Anforderung war, wurden diese Produkte aus der weiteren Betrachtung gestrichen. Unter den übrig gebliebenen Tools ist besonders das Tool Password Manager Pro [Man14] hervorgestochen, das die oben genannten Kriterien in fast allen Punkten unterstützt. Dieses Tool ist jedoch für einen Unternehmenseinsatz ausgelegt, bei dem wenige Administratoren vielen Nutzern gegenüberstehen. Im alltäglichen Umfeld eines Hochschul-Rechenzentrums ist dies jedoch nicht der Fall, so dass die Lizenzkosten für diese Lösung oberhalb des am LRZ vertretbaren Maßes lagen.

Aufgrund positiver Erfahrung einer Gruppe des LRZ und einer hausweiten Testphase wurde schließlich die Entscheidung für das Tool Password Safe Enterprise Edition (PSE) der Mateso GmbH [Mat14, vEH13] getroffen, um damit ein Notfall-Passwortmanagement am LRZ aufzubauen. Diese Software bietet eine zentrale Datenbank, die mit Hilfe von



Windows-Desktop-Clients angesprochen werden kann. Durch eine Synchronisation mit dem Active Directory ist die Authentifizierung einfach auf Mitarbeiter einschränkbar und durch die Möglichkeit der Versiegelung ist ein Notfall-Passwortmanagement mit dieser Software möglich. Zwar werden mit dieser Software einige der oben genannten Kriterien nicht (vollumfänglich) erfüllt, jedoch wurden diese Punkte als nicht betriebskritisch für eine Einführung gesehen.

Die Nachteile von PSE für den konkreten Einsatz am LRZ sind unter anderem:

- Das Ändern von Passwörtern auf dem Zielsystem wird nicht unterstützt. Dies bedeutet zwar einen erhöhten Mehraufwand für die Administratoren, erleichtert jedoch die Einführung eines Passwortmanagement-Tools, da die Administratoren keine Kontrolle aus der Hand geben müssen.
- Das Tool prüft nicht, ob Passwörter bereits verwendet wurden. Diese Einschränkung ist bisher am LRZ nicht von hoher Relevanz, da auch andere passwortverarbeitende Systeme dieser Einschränkung unterliegen.
- Falls auf ein Passwort im Notfall zugegriffen und dabei das angebrachte Siegel gebrochen wurde, ist der Zugriff auf das Passwort von Dritten gesperrt. Erst muss von einem berechtigten Benutzer das Siegel wieder hergestellt bzw. entfernt werden, bevor eine weitere Verwendung des Eintrags möglich ist.
- Eine temporäre Freigabe, wie oben gefordert, ist nur in dem Maß möglich, dass das wiederhergestellte Siegel erneut als Zugriffsbarriere besteht. Da es nicht möglich ist, Passwörter auf dem Zielsystem automatisch ändern zu lassen, ist es nicht möglich, weitere temporäre Einschränkungen einzubauen.

Was vor der Inbetriebnahme nur unzureichend bedacht wurde, ist, dass in der Einrichtungsphase schon vorhandene Passwörter geeignet in die Passwortmanagement-Lösung importiert werden müssen. Zwar bietet PSE, wie im Allgemeinen alle Lösungen, die Möglichkeit, dass Passwörter in einer einfachen Form wie CSV-Dateien importiert werden können, jedoch ist durch vielfältige Berechtigungsvergaben, Sichtbarkeitseinstellungen und Versiegelungsoptionen der manuelle Aufwand für jedes einzelne Passwort so hoch, dass ein einfaches Copy & Paste des Passworts in die Datenbank zeitlich nicht relevant ist. Der langfristige Vorteil liegt jedoch auf der Hand: Während man bei einer Offline-Lösung, wie beispielsweise den Passwortzetteln im Tresor, bei jeder Änderung des Passworts der organisatorische Aufwand nahezu gleich groß ist, ist bei der hier vorgestellten Lösung der initiale Aufwand zwar sehr hoch, wird jedoch nur das Passwort in einem Eintrag geändert, so bleiben alle anderen Eigenschaften wie eingestellte Berechtigungen bestehen und müssen nicht zwangsweise neu bearbeitet werden.

Somit wird den Administratoren ein praktisches Werkzeug in die Hand gegeben, das hilft, das tägliche Chaos mit Zugriffsberechtigungen zu bändigen. Die Akzeptanz des Programms unter den Administratoren entsprach jedoch zunächst nicht den hoch gesteckten Zielen während der Einführung. So haben zum aktuellen Zeitpunkt etwa 70 % aller Administratoren einen Zugriff auf das Passwortmanagement-Tool und es wird eine knapp vierstellige Anzahl von Passwörtern darin verwaltet. Zwar ist diese Zahl, verglichen mit

den am LRZ erbrachten Diensten noch recht bescheiden, aber aus Sicht der Administratoren wird die Datenbank im Wesentlichen für die kritischen Zugänge verwendet, die im Notfall bekannt sein müssen. Kleinere, unwichtige Systeme, deren Abschalten im Notfall keinen größeren Schaden verursacht, werden somit nicht über das Notfall-Passwortmanagement verwaltet sondern weiterhin über die dezentralen Lösungen. Die Gründe dafür sind sehr unterschiedlich. Neben einigen Unzulänglichkeiten des verwendeten Tools, wie beispielsweise einiger beschränkt aussagekräftiger Benachrichtigungen, die versendet werden, wenn ein versiegeltes Passwort angezeigt wird, und dem strategischen Fehler, in einer sehr heterogenen Umgebung auf eine damals rein Windows-basierte Lösung zu setzen, spielen die oben bereits angesprochenen Sicherheitsbedenken der Administratoren eine große Rolle.

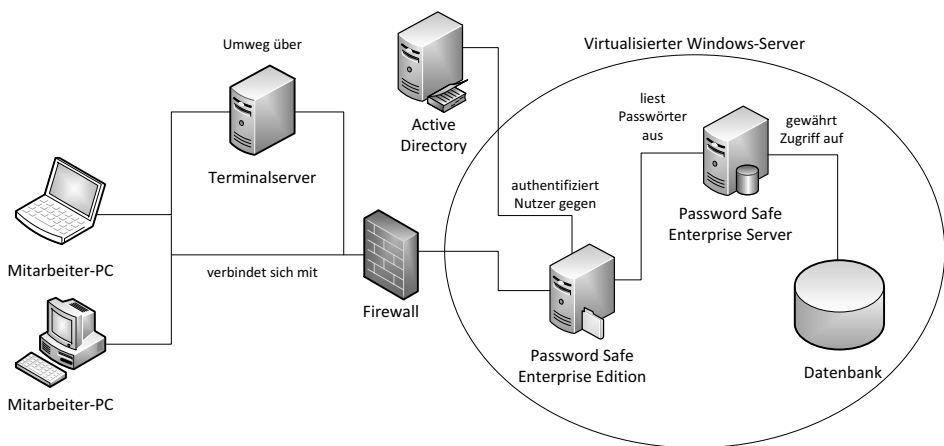


Abbildung 1: Zugriff auf die Notfall-Passwortmanagement-Lösung

Zwar wurde zur Absicherung der Passwort-Datenbank und des zentralen Servers eine Reihe von Maßnahmen ergriffen, wie beispielsweise die sehr strikte Einschränkung der Netz-basierten Erreichbarkeit des Servers, so ist, wie in Abbildung 1 zu sehen, im LRZ der Zugriff auf das Tool nur von Mitarbeiter-PCs oder über einen Terminalserver aus möglich, jedoch ist die generelle Ablehnung von Closed Source Software gerade im wissenschaftlichen Bereich auf einem hohen Niveau. Innerhalb der Datenbank wird jedoch in regelmäßigen Audits überprüft, ob die Administratoren die Berechtigungen und Sichtbarkeiten ihrer Passworteinträge korrekt gesetzt haben, so dass auf der einen Seite die Passwörter vor neugierigen Augen geschützt sind, auf der anderen Seite aber berechtigte Personen in Notfällen auf die Einträge zugreifen können.

Zusammenfassend kann man sagen, dass die Inbetriebnahme des Notfall-Passwortmanagements ein Erfolg war. Zwar sind auch nach gut einem Jahr Betrieb des Notfall-Passwortmanagements noch eine Reihe von Passwörtern nicht in das Tool eingetragen. Die Zahl der Passworteinträge ist jedoch kontinuierlich steigend und auch die Nutzung des Tools zur Erleichterung der täglichen Arbeit – da man sich das Merken oder anderweitig Notieren von Passwörtern sparen kann – wird von immer mehr Administratoren genutzt.

Diese Akzeptanzsteigerung wurde insbesondere auch durch begleitende Schulungsmaßnahmen, einer ausführlichen FAQ und nicht zuletzt durch eine persönliche Betreuung der Administratoren erreicht.

## 4 Fazit

Organisationsweit zentrales Notfall-Passwortmanagement hat zahlreiche Vorteile, stößt bei seiner Einführung aber häufig auf die Skepsis der Administratoren, die ihre Zugangsdaten darin hinterlegen sollen. Die Anzahl von Passwörtern, die mehreren Personen zugänglich sein müssen, kann zwar minimiert, aber in der Praxis nur sehr selten auf null gesenkt werden. Die Umstellung von rein persönlichem Passwortmanagement oder Offline-Verfahren wie verschlossenen Kuverts im Tresor auf ein organisationsweites Notfall-Passwortmanagement-Tool muss entsprechend gut geplant werden. Der größte Vorteil serverbasierter Lösungen ist die zuverlässige Alarmierung der zuständigen Administratoren, wenn sich jemand z. B. im Notfall einen temporären Zugriff auf einen Passworteintrag verschafft. Der Aufwand, beispielsweise die Ordnerstruktur, Default-Berechtigungen und E-Mail-Alarmierung zu konfigurieren, darf hierbei aber nicht unterschätzt werden.

**Danksagungen** Teile dieser Arbeit wurden durch das Bundesministerium für Bildung und Forschung gefördert (FKZ: 16BP12309). Die Autoren danken den Mitgliedern des Munich Network Management Teams (MNM-Team) für wertvolle Kommentare zu früheren Versionen dieses Artikels. Das MNM-Team ist eine Forschungsgruppe der Münchener Universitäten und des Leibniz-Rechenzentrums der Bayerischen Akademie der Wissenschaften unter der Leitung von Prof. Dr. Dieter Kranzlmüller und Prof. Dr. Heinz-Gerd Hegering.

## Literatur

- [BS12] Andrey Belenko und Dmitry Sklyarov. „Secure Password Managers“ and „Military-Grade Encryption“ on Smartphones: Oh, Really? *Blackhat Europe*, 2012.
- [Bun] Bundesamt für Sicherheit in der Informationstechnik. BSI-Grundsatzkataloge – Regelung des Passwortgebrauchs. <https://www.bsi.bund.de/ContentBSI/grundsatz/kataloge/m/m02/m02011.html>.
- [Bun11] Bundesamt für Sicherheit in der Informationstechnik. BSI gibt Tipps für sichere Passwörter. [https://www.bsi.bund.de/ContentBSI/Presse/Pressemitteilungen/Presse2011/Passwortsicherheit\\_27012011.html](https://www.bsi.bund.de/ContentBSI/Presse/Pressemitteilungen/Presse2011/Passwortsicherheit_27012011.html), 2011.
- [FH07] Dinei Florencio und Cormac Herley. A Large-scale Study of Web Password Habits. In *Proceedings of the 16th International Conference on World Wide Web, WWW '07*, Seiten 657–666, New York, NY, USA, 2007. ACM.

- [Gem12] Eva Gembler. *Erarbeitung eines kryptografischen Modells organisationsweiter Passwortverwaltung am Beispiel des Leibniz-Rechenzentrums*. Bachelorarbeit, Ludwig-Maximilians-Universität München, München, September 2012.
- [Man14] ManageEngine. Password Manager Pro. <http://www.manageengine.com/products/passwordmanagerpro/>, 2014.
- [Mat14] Mateso GmbH. Password Safe Enterprise Edition. <http://www.passwordsafe.de/>, 2014.
- [Mau08] Thomas Maus. Das Passwort ist tot – lang lebe das Passwort! *Datenschutz und Datensicherheit - DuD*, 32(8):537–542, 2008.
- [RZ06] Heiko Roßnagel und Jan Zibuschka. Single Sign On mit Signaturen. *Datenschutz und Datensicherheit - DuD*, 30(12):773–777, 2006.
- [vEH13] Felix von Eye und Wolfgang Hommel. Gut gesichert? – Anforderungen an ein zentralisiertes Passwortmanagement. In *ADMIN-Magazin*, Jgg. 03, Seiten 96–101, München, Deutschland, Mai 2013. Linux New Media AG.



# SIEGE: Service-Independent Enterprise-Grade protection against password scans

Marcel Waldvogel      Jürgen Kollek

Computer Science Department  
University of Konstanz  
Universitätsstr. 10, 78457 Konstanz, Germany  
<first>.<last>@uni-konstanz.de

**Abstract:** Security is one of the main challenges today, complicated significantly by the heterogeneous and open academic networks with thousands of different applications. Botnet-based brute-force password scans are a common security threat against the open academic networks. Common defenses are hard to maintain, error-prone and do not reliably discriminate between user error and coordinated attack. In this paper, we present a novel approach, which allows to secure many network services at once. By combining in-app tracking, local and global crowdsourcing, geographic information, and probabilistic user-bot distinction through differential password analysis, our PAM-based detection module can provide higher accuracy and faster blocking of botnets. In the future, we aim to make the mechanism even more generic and thus provide a distributed defense against one of the strongest threats against our infrastructure.

## 1 Introduction

Account information is in high demand among crooks. Hacked accounts are used (1) as stepping stones to hide their tracks, (2) to abuse legitimate web sites for phishing and distribution of malware and copyrighted material, (3) to send spam without being filtered, (4) for installing keyloggers to obtain more data, including banking information, and (5) to add machines to global botnets for later automated large-scale malfeasance.

One of the favorite uses of botnets is to farm more accounts to increase the power (and rental value) of the botnet [Pur03]. This is generally done by trying to authenticate against a service with (username, password) tuples based on dictionaries. These dictionaries are frequently enhanced by formation rules to e.g. append a few digits to words or replace the letter *e* by the digit 3.

In the beginning, these (username, password) pairs were fired at the target as quickly as possible. However, tools such as DenyHosts [Den] and Fail2ban [Fai] made it easy to defend against these types of attacks by blacklisting the IP address of the assailant.

Nowadays, the attacks happen much more subtly. A large set of bots, often several hundred thousand, jointly and covertly go stalking potential victims: Each of them only uses a few shy attempts, hoping this will not alert the subject. Later, another member of the hunting party will tickle the prey, while the first attacker approaches the next target.

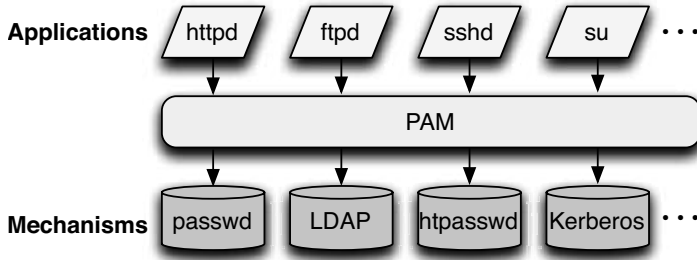


Figure 1: Flexibility of Pluggable Authentication Modules (PAM)

As it is necessary to allow a user to mistype her password or switch to the correct password for the system, a single failed attempt cannot be used as a direct indication of an attack. Identifying user-generated typos helps distinguishing between botnet attacks and failed user logins.

Our PAM[SS95]-based approach includes several firsts:

1. Retaining the power of PAM, allowing arbitrary authentication and authorization mechanisms to be used (Section 3)
2. Allowing geographical risk differentiation (GRID, Section 4.2)
3. Hierarchical risk grouping (Section 4.2)
4. Coordinated organisation-wide defense (CROWD, Section 4.1)
5. Dictionary-attack defense (Section 4.3)
6. Secure typo detection for user-bot differentiation (Section 4.4)

Based on flexible, realtime in-app tracking and combined with the use of DNS-based black lists, SIEGE provides low-maintenance, high-impact defense against password-guessing attacks.

## 2 Related work

Around 1990, it dawned to administrators and system programmers that traditional `/etc/passwd`-based login approach was not flexible enough. Kerberos [KN03], LDAP [WHK97] and derivatives such as Microsoft’s Active Directory started coming of age in the following decade. The Open Software Foundation came up with Pluggable Authentication Modules (PAM) [SS95] as the way of making authentication more flexible (Figure 1).

For each application, the system administrator can individually specify which authentication mechanisms to allow as part of the module chain (Figure 2 (b)). In addition to authentication, other account-specific functions, session management and change of authentication token (mostly passwords) can be controlled (Figure 2 (a)). Common functions include limiting access by date and time; creating and populating home directories on the first login; mounting special volumes such as encrypted homes (and securely unmounting

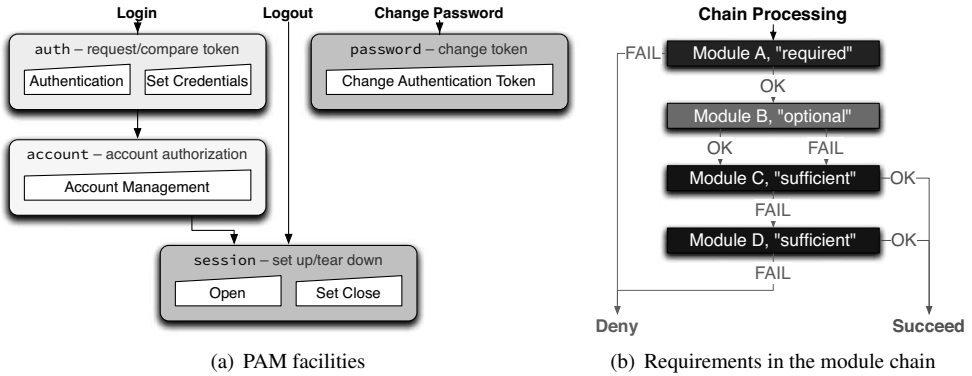


Figure 2: PAM architecture

them on logout); and many more (Figure 3).

Due to its versatility, PAM has rightly become the *de facto* authentication mechanism for most Unix-like desktop and server systems [Mor97].

For more than a decade, brute force account/password guessing attacks against SSH (and other services requiring authentication) have been going on [Cid13]. In the early days, a single host attacked its target with a large set of login attempts. This was simple to mitigate: Specialized host-based intrusion detection systems (IDS) like Fail2ban or DenyHosts [Fai, Den] started scanning log files for IP addresses which exceeded a given failed login rate threshold. The culprits were blocked from further attempts for a specified length of time.

Even though these are still not standard practice when setting up new hosts, the attackers moved away from individual source hosts to entire botnets. This helped them also to cover their tracks. Nowadays, each of the several ten or hundred thousands members of a botnet only probes a small set of username/password combinations before focusing on the next victim. This renders source-based limits essentially useless. Furthermore, log file analysis has several drawbacks: It must be flexible enough to cope with changes of the authentication software, but it must not be fooled into misinterpreting log entries.

For example, an `ssh` error message documenting a failed attempt for user `test` might look as follows

```
Jan 08 13:39:55 testhost sshd[555]: Invalid user test from 111.222.33.44 .
```

Malicious users can try to lock out users from a different address, say, 22.33.44.55, by trying to authenticate with the non-existent user `test` from 22.33.44.55 . The resulting message

```
Jan 08 13:39:59 testhost sshd[559]: Invalid user test from 22.33.44.55 from 111.222.33.44
```

and its relatives were misinterpreted by many earlier versions of most software packages, as they did not have the benefit of the differing background colors.

`pam_abl` [ABL] provides a similar function, but implements it as part of the PAM chain.



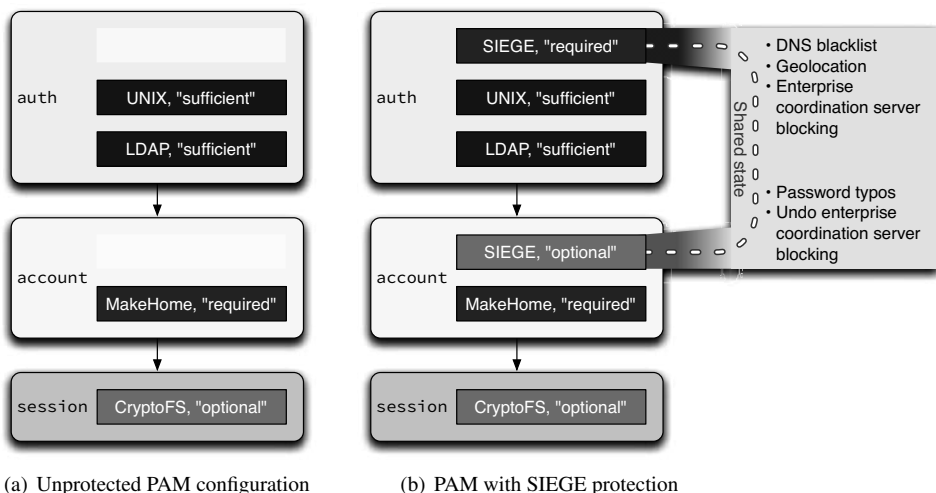


Figure 3: PAM protection

Therefore, it does not need to parse log files and thus is not vulnerable to the above log injection attacks. However, a host with `pam_abl` still fights alone in a steep uphill battle against a fierce horde of coordinated botnet members.

In more restricted environments (stronger rules, fewer services, or fewer users) than in academia, it is also possible to limit the hosts which may login or use certificate-based logins (including `ssh` keys). However, the diverse student and staff body of most universities makes the necessary education and support problems look like an unsurmountable problem, especially as this diverse group of ‘locals,’ together with many academic guests, make traditional security perimeters impossible: *inside* is the same as *outside*.

Intrusion Detection Systems (IDS) can try and limit logins, e.g. based on IP-based blacklists [Cis09]. However, due to the encrypted nature of `ssh` traffic, it is impossible to distinguish failed logins from short successful logins, which occur e.g. in remote management. Neither is it possible to differentiate between ‘good’ and ‘bad’ users at this level.

### 3 Operation

*SIEGE* brackets the actual authentication modules (Figure 3 (b)).<sup>1</sup> With this trick, it can remain mechanism-agnostic while learning as much as possible about successful and failed authentications: A successful authentication will be visible to both module parts, while a failed attempt will only call the top half.

<sup>1</sup>In a future version, we might consider bracketing it inside the authentication facility, between the *Authentication* and *Set Credentials* functions, similar to `pam_abl`.

The functions are divided into the top and bottom halves as follows.

### **Top half**

1. Check whether the remote host is known to third-party DNS-based blacklists, such as Spamhaus XBL [Spa]. Fail immediately on match.
2. Compare the password against a dictionary of common passwords (Section 4.3)
3. Determine remaining number of failed attempts for this address:
  - (a) Look up source address in the CROWD (CooRdinated Organization-WiDe, Section 4.1) server. If unreachable, use local cache. Return this information, if any. Otherwise:
  - (b) Determine geographic origin of the connection. Use administrator-set likelihood for legitimate login attempts modified with optional user preferences to determine maximum number of login attempts (Section 4.2).
4. Decrease the number of remaining login attempts, indicating that a (so far unsuccessful) attempt is in progress. The amount of reduction is tuned by the likelihood of a user password mistake (Section 4.4).
5. Perform hierarchical blocking propagation (Section 4.2).
6. Store this information in the local cache and update the CROWD.

To simplify the process of writing back local results if the CROWD has been temporarily unavailable, the minimum of the remaining attempts of CROWD and the local cache can be used.

### **Bottom half**

1. Undo the attempted reduction and update cache and CROWD accordingly.

This unique three-level approach of (1) global, long-term DNS blacklist, (2) organization-wide CROWD, and (3) local system information, is able to strike an excellent balance between security, manageability, reliability, control, and privacy. The local level can react quickly and provides defense even when the higher levels are unreachable, e.g. because of denial of service, while CROWD helps to react even to slow attacks. The DNS blacklists help identify the long-time perpetrators, even if they have been inactive for some days. These functions cannot all be bundled into one global system, as this itself would become a security risk. In our design, the global level never sees username information, and password-related information does not leave the local host and even there is only stored for a very short period of time.

Details about the implementation can be found in [Kol13].

## 4 Defense mechanisms

The challenge in the defense against password-guessing attacks is finding a balance between

1. locking out as many of the attackers as possible, to reduce their guessing rate as much as possible, and
2. allowing legitimate users to log in even when an attack is ongoing.

Our approach at achieving this is described in the remainder of this section.

### 4.1 Coordinated organization-wide defense

In the spirit of Indra [JWZ03], CROWD uses the wisdom of the (cooperating) masses within an organization to coordinate their defense against the seemingly unsurpassable superiority of the botnet. Each host contains a local secret, which is used to authenticate against the central CROWD server through a secure connection.

The race conditions in the protocol are hard to exploit, as the problem windows are very short-lived. This approach still allows the clients to remain in full control, keep the server simple (get, put, expire) and close gaping accounting loopholes in other approaches. For example, `pam_abl` (and, depending on the configuration, also the other mechanisms), allow a single-source attacker to open a large number of connections simultaneously. When they are open and ready to receive the password, hundreds or thousands of parallel guesses can be made, before the source is blacklisted. Due to the pre-decrementing used in SIEGE, even with some decrements missing due to the race condition, only the first few attempts will be evaluated before the source is stopped.

### 4.2 Geographical risk differentiation

Even at a university with global outreach, the services requiring authentication have a local user base: Most requests will be from within the campus, the city, state, or country. Therefore, remote locations are less likely to be the source of legitimate logins, making password mismatches from other locations more likely to be from attackers. As a result, the threshold to block a given source for exceeding the number of bad password attempts can be set smaller while maintaining a high probability that legitimate users will not be locked out.

SIEGE uses this property to differentiate failure thresholds and their effects. A sample set of parameters is given in Table 1: For the university's home country, after 10 failed attempts within a defined period (typically 10 minutes), the host is blocked for a configurable time (10 minutes as well). As soon as 10 hosts within the same subnet (/24 in the case of IPv4) are blocked, the entire subnet is blocked for 20 minutes. When the blocked subnets in a /16

| Class     | Failure thresholds |        |     |          | Blocking periods (minutes) |        |     |         |
|-----------|--------------------|--------|-----|----------|----------------------------|--------|-----|---------|
|           | Host               | Subnet | Net | Country  | Host                       | Subnet | Net | Country |
| Home      | 10                 | 10     | 10  | $\infty$ | 10                         | 20     | 30  | —       |
| Neighbors | 5                  | 5      | 5   | 20       | 10                         | 20     | 30  | 60      |
| Others    | 2                  | 2      | 2   | 10       | 10                         | 20     | 30  | 60      |

Table 1: Example escalation parameters for password failures

network reaches 10, the entire /16 network is blocked for 30 minutes. Even when many networks in the home country are blocked, the country itself is never completely blocked. For neighbor and other countries, entire countries can be blocked.

This mechanism can be too aggressive for a member of the university currently visiting a country where a large number of bots are currently operating. Therefore, we allow individual users to list countries they are visiting.<sup>2</sup> For those users, the login threshold will be set to 2 (or the minimum of the country) and they will be exempt from network size escalation.

This allows to trade the small risk of that user’s account being attacked against the convenience of the user. Because this has to be a user-driven step, the user can also be educated about security precautions at this time.

With this adaptive combination of reactions, the GRID component of SIEGE can ensure a very high level of security without compromising user experience.

### 4.3 Dictionary-based botnet identification

As a rule of thumb (and often as part of the security policy), passwords should not be derived from dictionary words or well-known frequently-used passwords. This can be enforced by the CrackLib PAM module [Gaf] or other similar tools. Conversely, this also means that a user is unlikely to use a dictionary word as their password. SIEGE includes a simple test whether a password tried is known to CrackLib and immediately classifies a host trying to use such a word as an attacker, putting it on a temporary blacklist. This significantly increases the efficiency of the defense perimeter.

### 4.4 User-/bot password differentiation

For processes supporting manual password entry, a user has a chance of mistyping the password. Common typing mistakes include:

**Transposing** two keys on the keyboard: Two adjacent characters in the password are swapped.

---

<sup>2</sup>This is best done before leaving the home country, to avoid a lockout in the case of an acute attack.

**Duplicating.** Bouncing a key on the keyboard: A character is in the password twice instead of once.

**Missing.** Not hitting a key strong enough: The resulting password is missing a character.

It is unlikely that an attacker will try the correct password modified by one of these rules early in the process. Thus, seeing a password modified by one of these rules is likely to come from a user. When we detect this, the try will not count as much as a full incorrect password attempt, thus reducing the chance of a legitimate user being locked out incorrectly.

As the above error-generating processes can all be reversed, this is easy, even when the passwords are encrypted (which they should be):

**Transposing.** All possible single transpositions can be tried in  $L - 1$  password encryption attempts locally on the machine, where  $L$  is the length of the password just received.

**Duplicating.** All sequences of duplicate successive characters can be eliminated each in  $D$  password encryption attempts locally on the machine, where  $D$  is the number of duplicate successive characters ( $D < L$ ).

**Missing.** All possible missing characters can be tried in  $(L + 1) * A$  password encryption attempts locally on the machine, where  $A$  is the size of the likely alphabet.

The first two mechanisms are very cheap, the third can be rather costly if an expensive encryption function was chosen for the password, so in some environments, this costly step might be undesirable and can thus be deactivated.

Mechanisms which perform weak matches against the password open up new attack vectors. One prominent mechanism are timing attacks, where a shorter password verification time would indicate that the correct password has been found after only a subset of the operations. To thwart such timing attacks, all the configured operations must be performed for an incorrect input password. As the number of password verification options then only depend on the input password, the timing leaks no information about the proximity of the input to the actual password.

As the weak matching does not directly return information whether the password was a close match or not, this information is very hard to abuse. Therefore, we believe that this option only strengthens the defense, not weakens it, as it allows for significantly narrower limits before blacklisting.

## 5 Evaluation

We built a simulator of the GRID and hierarchical blocking propagation logic to show the large-scale effects (Table 2). For this simulation, we assumed three botnet sizes, ranging from 10 thousand to 100 million hosts. The botnets are following either of two goals: A targeted attack at a single host, or a sweeping attack, which ‘just’ tries to collect accounts, which has identified 1,000 hosts on the campus. We also estimate that due to protocol overhead and built-in backoff timers, each bot can try a single password per second per

| Target hosts       | single |             |             | 1k      |             |              |
|--------------------|--------|-------------|-------------|---------|-------------|--------------|
| Botnet hosts       | 10k    | 1,000k      | 100,000k    | 10k     | 1,000k      | 100,000k     |
| Unprotected        | 10k    | 1,000k      | 100,000k    | 10,000k | 1,000,000k  | 100,000,000k |
| Fail2ban/DenyHosts | 0.2k   | 16k         | 1,630k      | 163k    | 16,300k     | 1,630,000k   |
| SIEGE, 1 home      | 0.0k   | 0.1k        | 5k          | 0.0k    | 0.1k        | 5k           |
| SIEGE, 5 neighbors | 0.0k   | 0.2k        | <b>0.5k</b> | 0.0k    | 0.2k        | <b>0.5k</b>  |
| SIEGE, 250 others  | 0.0k   | <b>0.4k</b> | <b>0.4k</b> | 0.0k    | <b>0.4k</b> | <b>0.4k</b>  |
| SIEGE, total       | 0.0k   | 0.7k        | 6.5k        | 0.0k    | 0.7k        | 6.5k         |

Table 2: Password trial rate [ $s^{-1}$ ] (bold numbers indicate country blocking)

target host.

For the unprotected hosts, this means that the total number of password attempts is the product of botnet size and target hosts, resulting in up to stunning  $10^{11}$  password attempts per second!

Fail2ban and DenyHosts are assumed to block a host after 10 failures for the following 10 minutes, so the effective trial rate is 10 attempts every 610 s, improving on the unprotected hosts by a factor of 61.

SIEGE uses the parameters from Table 1. We assume that there are 256 countries, one home country, five neighbors, and the remaining 250 in the ‘other’ class. Each country consists of  $256/16$  nets, subdivided into  $/24$ s. Attacking hosts are uniformly distributed about the resulting  $2^{32}$  addresses. This is actually the worst case for SIEGE, as any clustering will greatly increase the impact of the hierarchical blocking. We also assume that the attackers will carefully avoid all passwords found in the targets’ dictionary, to avoid multiplying SIEGE blocking rate (and thus significantly reducing attack rate).

The results are impressive: SIEGE limits background-type scans to a few dozen password attempts per second. But even under heavy attack, no more than 7k attempts are successful, while still offering service to most legitimate users. GRID and CROWD combined offer 7 orders of magnitude protection over vanilla systems and more than 5 magnitudes compared to Fail2ban/DenyHosts-style.

## 6 Conclusions

Even a small campus without any specific protection can be subject to roughly  $10^{17}$  password scans per day, especially on the weekend, when monitoring is not as close. This is equivalent to trying *all(!)* twelve-digit lowercase passwords, ten-digit alphabetic passwords or nine-digit passwords perfectly chosen from all 95 ASCII printable characters. With a large number of unprotected hosts, you no longer should assume that even excellent passwords are safe. Therefore, alternate mechanisms are needed. GRID and CROWD mechanisms of SIEGE provide several orders of magnitude improvement to current approaches. When you know that your campus is free of dictionary-derived passwords, the attacker identification can be significantly sped up, providing even stronger protection.

We are currently working on putting the code into a releasable format to make the SIEGE benefits available to the community.

## References

- [ABL] PAM ABL: Automatic defense against brute force attacks. <http://pam-abl.sourceforge.net>, accessed 2014-04-08.
- [Cid13] Daniel Cid. SSH brute force – The 10 year old attack that still persists. Sucuri Blog, July 2013. <http://blog.sucuri.net/2013/07/ssh-brute-force-the-10-year-old-attack-that-still-persists.html>, accessed 2014-04-08.
- [Cis09] Combating Botnets Using the Cisco ASA Botnet Traffic Filter. white paper, Cisco, 2009. [http://www.cisco.com/c/en/us/products/collateral/security/asa-5500-series-next-generation-firewalls/white\\_paper\\_c11-532091.pdf](http://www.cisco.com/c/en/us/products/collateral/security/asa-5500-series-next-generation-firewalls/white_paper_c11-532091.pdf).
- [Den] DenyHosts. <http://denyhosts.sourceforge.net>.
- [Fai] Fail2ban. <http://www.fail2ban.org>.
- [Gaf] Cristian Gafton. pam-cracklib - checks the password against dictionary words. [http://www.linux-pam.org/Linux-PAM-html/sag-pam\\_cracklib.html](http://www.linux-pam.org/Linux-PAM-html/sag-pam_cracklib.html), accessed 2014-04-08.
- [JWZ03] Ramaprabhu Janakiraman, Marcel Waldvogel, and Qi Zhang. Indra: A Peer-to-Peer Approach to Network Intrusion Detection and Prevention. In *Proceedings of IEEE WETICE 2003*, Linz, Austria, 2003.
- [KN03] J. Kohl and C. Neumann. The Kerberos Network Authentication Service (V5). RFC 1510, Internet Engineering Task Force (IETF), September 2003. <http://tools.ietf.org/html/rfc1510>.
- [Kol13] Jürgen Kollek. Ein verbesserter PAM-basierter Ansatz um “brute force”-Passwort-Angriffe auf den “secure shell”-Service zu erkennen und zu verhindern. Bachelor’s thesis, University of Konstanz, 2013. <http://kops.ub.uni-konstanz.de/handle/urn:nbn:de:bsz:352-259480>.
- [Mor97] Andrew G. Morgan. Pluggable Authentication Modules for Linux. *Linux Journal*, December 1997.
- [Pur03] Ramneek Puri. Bots & Botnet: An Overview. SANS Institute InfoSec Reading Room, August 2003. <http://www.sans.org/reading-room/whitepapers/malicious/bots-botnet-overview-1299>, accessed 2014-04-08.
- [Spa] Spamhaus. XBL – Exploit and Botnet filter. <http://www.spamhaus.org/xbl/>.
- [SS95] V. Samar and R. Schemers. Unified Login with Pluggable Authentication Modules (PAM). Request for Comments 86.0, Open Software Foundation, October 1995. Archived at <http://www.kernel.org/pub/linux/libs/pam/pre/doc/rfc86.0.txt.gz>, accessed 2014-04-08.
- [WHK97] M. Wahl, T. Howes, and S. Kille. Lightweight Directory Access Protocol (v3). RFC 2251, Internet Engineering Task Force (IETF), December 1997. <http://tools.ietf.org/html/rfc2251>.

# Architektur zur mehrstufigen Angriffserkennung in Hochgeschwindigkeits-Backbone-Netzen

Mario Golling, Robert Koch und Lars Stiemert  
Munich Network Management Team (MNM-Team)  
Universität der Bundeswehr München  
Werner-Heisenberg-Weg 39  
D-85577 Neubiberg  
{mario.golling, robert.koch, lars.stiemert}@unibw.de

**Abstract:** Die globale Vernetzung und die Durchdringung des alltäglichen Lebens durch Informations- und Kommunikationstechnologien, sowie die zunehmende Anzahl von Angriffen, die bisweilen auch unter Beteiligung ahnungsloser Nutzer durchgeführt werden (Bot-Netze) führen dazu, dass Angriffe auf IT-Infrastrukturen längst keine zu vernachlässigende Begleiterscheinung des Internets mehr sind. Angriffe wie der auf das Spamhaus Project, das 2013 durch einen Distributed Denial of Service Angriff mit mehr als 300 Gbps attackiert wurde und in der Konsequenz auch einige Backbone Provider an die Grenze Ihrer Leistungsfähigkeit brachte, zeigen eindrücklich, dass Angriffe auch auf Backbone Provider große Auswirkungen haben können. Systeme zur Angriffserkennung arbeiten historisch betrachtet zumeist auf Basis sogenannter Deep Packet Inspection, bei welcher der Paketinhalt auf das Vorhandensein spezieller Muster überprüft wird. Dies ermöglicht zwar detaillierte Analysen, ist aber aufgrund der mangelnden Skalierbarkeit in Zeiten immer schneller werdender Anbindungen, gerade für Backbone-Provider, finanziell wie auch technisch nicht praktikabel. Spezifische rechtliche Beschränkungen verschärfen diese Problematik zusätzlich. Die vorliegende Publikation stellt deswegen einen mehrstufigen Ansatz zur Angriffserkennung speziell für Backbone Provider vor, bei welchem mehrere Verfahren wie unter anderem Flow-basierte Angriffserkennung, Protokoll-basierte Angriffserkennung, Deep Packet Inspection und Geolokalisation kombiniert werden.

## 1 Einleitung und Problemstellung

Die Informations- und Kommunikationstechnik (IKT) hat sich zu einem wesentlichen Bestandteil des alltäglichen Lebens entwickelt. Die globale Vernetzung und die stetig voranschreitende Durchdringung des privaten wie öffentlichen Sektors eröffnet ungeahnte Möglichkeiten und fördert Innovationen ebenso wie die Wirtschaft als Ganzes. Einhergehend mit dem Bedeutungszuwachs des Internets nehmen allerdings auch die Anzahl der Angriffe in den letzten Jahren stetig zu - sowohl in quantitativer, als auch in qualitativer Ausprägung. Um den zunehmenden Bedrohungen entgegenzuwirken haben sich zur Detektion von Angriffen zahlreiche Lösungen etabliert, wie insbesondere Systeme



zur Angriffserkennung (*engl. Intrusion Detection Systeme; IDSs*). Aktuelle Systeme arbeiten hier vorwiegend nach dem Prinzip der Deep Packet Inspektion (DPI), bei der eine Angriffserkennung durch Analyse des kompletten Paketinhaltes (auf das Vorhandensein typischer Indikatoren) stattfindet. Während sich DPI-basierte Systeme aufgrund ihrer im Vergleich zu verhaltensbasierten Systemen besseren Fehlalarmraten, sowie der detaillierten Untersuchungsmöglichkeiten des Datenverkehrs als Standard-Sicherheitssysteme in Access Netzen etabliert haben, ist ihr Einsatz in Netzen mit höheren Verbindungsgeschwindigkeiten (größer 10 Gigabit) mit sehr hohen Kosten verbunden bzw. technisch schwer realisierbar. Zusätzlich bestehen bei DPI zahlreiche rechtliche Bedenken hinsichtlich der Privatsphäre und des Datenschutzes, ebenso wie der juristischen Zuständigkeit.

Insbesondere wenn Angriffe verteilt durchgeführt werden, kann ihr Gesamtvolumen die jeweiligen Systeme überlasten. So wurde zum Beispiel das Spamhaus-Projekt von Distributed Denial of Service (DDoS)-Angriffen Anfang 2013 mit mehr als 300 Gbps Verkehr belegt - genug, um einzelne Knotenpunkte des Internets zu überlasten. Zusätzlich zu DDoS-Attacken stellen v.a. Würmer und Botnetze besondere Herausforderungen für Backbone-Provider dar, da auch sie dazu neigen, eine große Menge an Ressourcen zu blockieren [Arb]. Im Rahmen dieser Publikation wird der Begriff Backbone Provider verwendet, um Netze mit Geschwindigkeiten höher als 10 Gbps und ohne eine Verbindung zu spezifischen Endsystemen zu bezeichnen. Die Kenntnis des geographischen Ursprungs eines Angriffes kann heutzutage von großer Bedeutung sein. So hat der Sicherheitsdienstleister Mandiant kürzlich im Rahmen einer Studie festgestellt, dass ein großer Teil aller Angriffe auf Behörden der USA auf den näheren Umkreis eines einzigen Gebäudes in Shanghai/China einschränkbar ist [Man13].

Dieser Problematik Rechnung tragend wird im Rahmen dieser Publikation ein mehrstufiger Ansatz zur Angriffserkennung speziell für Backbone-Provider vorgestellt, indem mehrere Verfahren wie unter anderem Flow-basierte Angriffserkennung, Protokoll-basierte Angriffserkennung, DPI und Geolokalisation kombiniert werden.

Hierzu wird in Abschnitt 2 ein Überblick über vorhandene Verfahren - sowohl im Bereich der Angriffserkennung, als auch im Bereich der Geolokalisation gegeben. Anschließend wird in den Abschnitten 3 und 4 die mehrstufige Architektur sowie die prototypische Implementierung vorgestellt. Im letzten Teil dieser Arbeit erfolgt ein kurzes Resümee sowie ein Ausblick auf zukünftige Arbeiten.

## **2 Stand der Wissenschaft und Technik**

Im folgenden Abschnitt erfolgt zunächst ein kurzer Überblick über Verfahren zur Angriffserkennung. Analog dazu werden danach im zweiten Teil vorhandene Ansätze zur Geolokalisation vorgestellt und es wird kurz darauf eingegangen, inwiefern diese vorteilhaft in den Prozess der Angriffserkennung integriert werden können.

## 2.1 Angriffserkennung

Gemäß des National Institute of Standards and Technology (NIST) wird Angriffserkennung respektive Intrusion Detection allgemein als Prozess zur Überwachung und Analyse von Ereignissen in Netzen oder Computersystemen verstanden, bei dem Vorkommnisse auf abweichende Verhaltensweisen hinsichtlich der Verletzung von Sicherheits-, Nutzungsrichtlinien oder -praktiken untersucht werden [SM10].

### 2.1.1 Vorhandene Ansätze zur Angriffserkennung

Im Bereich der Angriffserkennung haben sich über die Jahre verschiedene Formen von IDSs herausgebildet. Eine der grundlegenden Unterscheidungen hierbei ist die Differenzierung in hostbasierte Systeme und solche Systeme, bei denen die Erkennung nicht auf dem Computersystem erfolgt, sondern mittels Sensoren im Netz. Netzbasierte IDS (NIDS) lassen sich ebenfalls noch weiter unterteilen und werden im Folgenden basierend auf der ISO/OSI-Ebene, auf der die Angriffserkennung durchgeführt wird, kurz vorgestellt:

**Angriffserkennung auf Basis der Auswertung des Paketinhaltes:** Der bereits vorgestellte DPI-Ansatz stellt den gebräuchlichsten Vertreter von NIDSs dar. Charakteristisch für DPI ist, dass hier der gesamte Inhalt jedes Paket auf das Vorhandensein von Auffälligkeiten untersucht wird. In der Praxis erfolgt hierbei im Regelfall ein Abgleich der zu untersuchenden Datenpakete zu vorher gewonnenen Signaturen.

**Angriffserkennung unter Auswertung statistischer, bzw. protokollspezifischer Informationen:** Neben der DPI-basierten Angriffserkennung auf Schicht 7 des ISO/OSI-Referenzmodells kann aber auch eine Auswertung auf Schicht 4, d.h. eine Auswertung des Nachrichtenkopfs (*engl. Header*) erfolgen. Aufgrund der Tatsache, dass nur Header untersucht werden, sind diese Ansätze auch in der Lage, hohe Datenraten zu verarbeiten. Hierbei zeichnen sich *Protokoll-basierte IDSs* dadurch aus, dass sie das dynamische Verhalten und den Zustand des Protokolls bzgl. des durch den Standard vorgesehenen Ablaufs hin untersuchen. Problematisch hierbei ist, dass Standards oftmals nicht alle Aspekte des Protokollablaufs festlegen (bspw. Start von Sequenznummern - zufallsbasierter Anfangswert oder Null, etc.) und Hersteller sich bei der Implementierung nicht immer komplett an den Standard halten; dies hat maßgeblichen Einfluss auf die Fehlalarmraten eines entsprechenden IDS' und kann nur durch die Untersuchung und Berücksichtigung der spezifischen Abweichungen verbessert werden. *Verhaltensbasierte IDSs* hingegen beruhen auf Modellen wie bspw. dem Satz von Bayes, um bösartige Pakete im Netz zu identifizieren. Als Folge sind statistische IDS in der Lage, ihr Systemverhalten dynamisch anzupassen und damit ihre eigene Regeln zu modifizieren. Statistische Auswertungen sind grundsätzlich sowohl auf Schicht 4, als auch auf Schicht 7 des ISO/OSI-Modell möglich, wobei im Zuge dieser Arbeit der Begriff nur für Schicht 4-basierte Systeme Verwendung findet.

**Angriffserkennung auf Basis aggregierter Header-Informationen:** Anstelle der Auswertung von *allen* Headern kann eine Angriffserkennung auch auf Basis aggregierter Verkehrsdaten (Flows) erfolgen [SSS<sup>+</sup>10]. RFC 7011 [CTA13] definiert einen Flow als Menge von IP Paketen mit gemeinsamen Eigenschaften, die innerhalb eines bestimmten Zeitrahmens erfasst werden. Diese Art der Auswertung ermöglicht eine Analyse, welche rechtliche Bedenken weitgehend ausräumt, da die eigentlichen Paketinhalte nicht herangezogen werden. Weiterhin ist aufgrund der reduzierten Datenmenge der Flows eine sehr effiziente Auswertung möglich; andererseits müssen die Flows zu den Paketströmen erst generiert werden, was eine Verzögerung zu anderen Detektionsverfahren mit sich bringt. Während bei DPI die Inhalte sofort betrachtet werden, muss für die Generierung der Flow-Daten erst eine bestimmte Anzahl an Paketen übertragen und berücksichtigt werden.

**Bewertung und Eignung für Backbone-Provider:** Für Backbone-Provider ist eine Angriffserkennung basierend auf Flows derzeit sowohl rechtlich, als auch vom Aufwand her betrachtet die einzig praktikable Möglichkeit. Viele Backbone-Provider setzen ohnehin Flow-basierte Komponenten im Rahmen der Netzüberwachung und des -managements ein. Diese Datenquellen können direkt für die Flow-basierte Angriffserkennung mit genutzt werden. Flow-basierte Angriffserkennung hat andererseits den Nachteil, dass nicht alle Arten von Angriffen und böartigem Verhalten erkannt werden können. Folgende vier Kategorien können gut erkannt werden: (i) (Distributed) Denial of Service DDoS (ii) Netz-Scans (iii) Würmer sowie (iv) Bot-Netze [SSS<sup>+</sup>10]. Da dies jedoch genau die Gefahren sind, die einen Backbone-Provider interessieren (bspw. die Detektion eines DoS und das Ergreifen von Gegenmaßnahmen) [Arb], ist eine Flow-basierte Detektion als ausreichend zu bezeichnen.

## 2.2 Geolokalisation

Geolokalisation bezeichnet die Zuordnung einer logische Adresse, beispielsweise der IP eines Hosts, zu einer physikalischen respektive geografischen Position. Das Einsatzspektrum ist hierbei breit gefächert und umfasst neben zielgerichteter Werbung beziehungsweise Contentpräsentation auch die Verwendung im Rahmen der Strafverfolgung, wenn es um die juristische Zuständigkeit oder den Ursprung eines Deliktes mit dem Tatmittel Internet geht.

### 2.2.1 Vorhandene Ansätze zur Geolokalisation

Im Forschungsgebiet der Geolokalisation von IP-Adressen haben sich diverse Ansätze herausgebildet. Nach Endo et. al [ES10] lassen sich alle Verfahren grundsätzlich in zwei wesentliche Kategorien einordnen [KGR13b]:

**Verfahren basierend auf aktiven Messungen:** Die Ermittlung der möglichen geografischen Position eines Hosts erfolgt über aktive Messungen. Es findet folglich eine direkte

Interaktion mit dem Zielsystem statt. Dazu werden typischerweise Verzögerungswerte (zumeist Round Trip Times) bestimmt und daraufhin versucht, über Vergleiche mit Messungen bekannter Server-Standorte (beispielsweise Webserver oder Router), auf die Position des Zielhosts zu schließen [ZFdRD05].

**Semantische Methoden:** Anhand semantischer Verfahren werden standortbezogene Informationen zu einer gegebenen Adresse mit Hilfe umfangreicher Datenanalysen ermittelt. Die so erhobenen Daten können anschließend Rückschlüsse auf die tatsächliche physikalische Position des Ziels zulassen. Grundlage hierfür bilden unter anderem die Datenbestände der fünf Regional Internet Registries, sowie Geoservice-Anbieter wie Quova [HFc11]. Diese Ansätze ermöglichen Abfragen in Echtzeit und sind entsprechend für große Datenmengen geeignet. Allerdings sind die Datenbestände zumeist nicht aktuell und folglich ungenau.

**Bewertung und Eignung für Backbone-Provider:** Im Rahmen der im nächsten Abschnitt vorgestellten Architektur hat Geolokalisation eine hohe Bedeutung. Da, wie bereits erwähnt, einer Studie des Sicherheitsdienstleister Mandiant zufolge ein großer Teil aller Angriffe auf Behörden der USA auf den näheren Umkreis eines einzigen Gebäudes in Shanghai/China einschränkbar ist [Man13], soll Geolokalisation vergleichbar mit Greylisting in E-Mails verwendet werden, um so verdächtigen Pakete (deren Ursprung sich in räumlicher Nähe zu ähnlichen, bereits erkannten, bösartigen Paketen befindet) einen höheren Angriffswahrscheinlichkeit zuzuordnen. Da die rechtlichen Rahmenbedingungen in der Regel länderspezifisch sind, ist die Nutzung von Geolokalisationsverfahren für die geographische Zuordnung eines Einbruchversuches zudem ein probates Mittel, um entsprechende (Gegen-)Maßnahmen auch vor länderspezifischen rechtlichen Besonderheiten gegeneinander abzugrenzen [KGR13a].

Bei Betrachtung der finanziellen und technischen Herausforderungen in heutigen Hochgeschwindigkeitsnetzen sind Geolokalisationsverfahren auf Basis von aktiven Messungen vor allem aufgrund des hohen Datenaufkommens nicht praktikabel einsetzbar. Obwohl widersprüchliche Angaben zur Genauigkeit von derlei Ansätze existieren, gehen jüngste Studien [HFc11] davon aus, dass eine Zuordnung auf Landesebene zu 96-98% korrekt ist. Als erweiternde Maßnahme zur Angriffserkennung, um beispielsweise den geographischen Ursprung von Angriffen zu korrelieren, liefern Geo-Datenbanken somit einen kostengünstigen und effektiven Mehrwert.

### 3 Architektur

Nachfolgend wird auf Basis der Erkenntnisse aus Abschnitt 2 ein mehrstufiger Ansatz zur Angriffserkennung vorgestellt (siehe Abbildung 1), welcher die Vorteile verschiedener Systeme zur Angriffserkennung wie DPI und Flow-basierter IDSs vereint, um so die Fehlalarmraten zu minimieren.

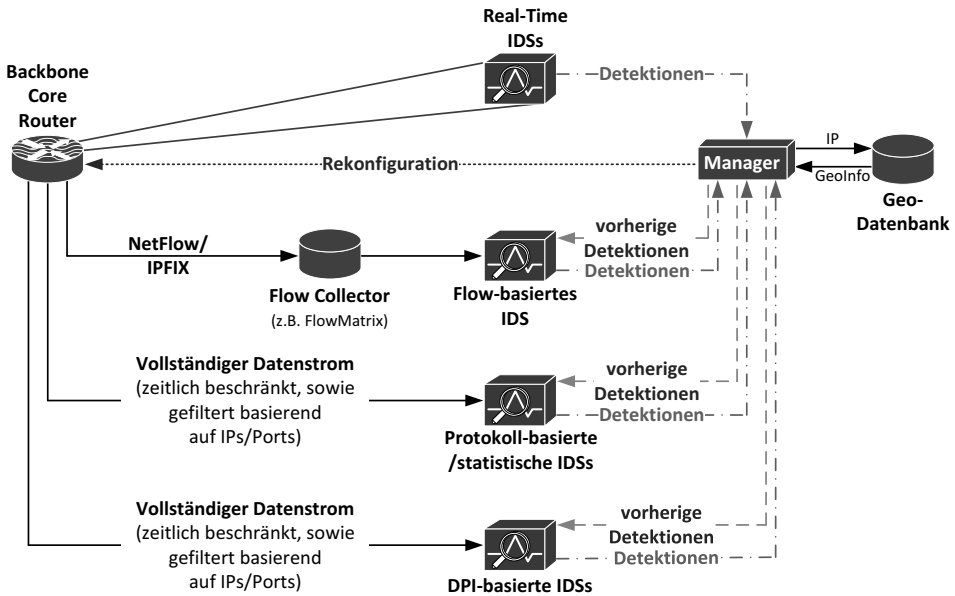


Abbildung 1: Komponenten der Architektur

**Grundlegendes Prinzip:** Wie in Abbildung 1 dargestellt, besteht die Architektur im Wesentlichen aus einer *Managementkomponente (Manager)*, einem *Core Router* und verschiedenen *IDSs*. Die Managementkomponente ist hierbei das Kernelement der Architektur. Sie steuert alle Interaktionen, kontrolliert die einzelnen IDSs und stellt sicher, dass diese bei Bedarf die nötigen Verkehrsdaten erhalten.

**Real-Time IDS:** Im Ausgangszustand wird dabei der komplette Datenstrom, der den Core Router passiert, durch ein spezielles Real-Time IDS, welches direkt auf dem Router ausgeführt wird, untersucht. Da hier kein dediziertes Gerät vorhanden ist, sondern dem Router Ressourcen „entzogen“ werden, sind nur sehr einfache Funktionen möglich. Die zentrale Idee hierbei ist, insbesondere gegen (D)DoS-Attacken zu schützen. So wurde bereits gezeigt, dass es sehr gut möglich ist, einen solchen Angriff mit den begrenzten Ressourcen eines Routers mit einer hohen Genauigkeit (selbst bei einer relativ hohen CPU-Auslastung) [HBSP13] zu erkennen. Im Falle des Erkennens eines Angriffes ist das Real-Time IDS in der Lage, selbständig den Router zu rekonfigurieren, um den Datenstrom bspw. durch eine Access Control List oder ein Sinkhole zu blockieren. Darüber hinaus wird auch der Manager unter Nutzung des „Intrusion Detection Message Exchange Formats“ (IDMEF, Austauschformat) sowie der „Common Vulnerabilities and Exposures“ (CVE, standardisierte Semantik) und dem „Common Vulnerability Scoring System“ (CVSS, standardisierte Metrik) über den Vorfall informiert.

**Flow-basierte IDS:** Neben dem Real-Time IDS überwacht auch das Flow-basierte IDS ständig den kompletten Datenstrom des Routers. Angesichts der Tatsache, dass Flow-

Exportformate wie NetFlow und IPFIX die Verkehrsdaten in Flows aggregieren, ist eine derartige Angriffserkennung auch mit vergleichsweise geringen Ressourcen möglich. Zur einfacheren Detektion kann das Flow-basierte IDS zudem vom Manager über bereits vorhandene Detektionsergebnisse informiert werden. Obwohl dies von gegenwärtigen Systemen nicht unterstützt wird, besteht die Grundidee darin, einem IDS so viele Information wie möglich an die Hand zu geben, um so den Prozess der Angriffserkennung so zuverlässig wie möglich zu machen. Sollte das Flow-basierte IDS einen Angriff feststellen, so leitet es die Information an den Manager weiter. Dieser wertet die Information aus und kann seinerseits wiederum andere IDSs (wie protokoll-basierte oder DPI-basierte IDSs) aktivieren und über die vorangegangene Detektion beziehungsweise die Meldungen der jeweiligen IDSs entsprechend steuern und informieren. In diesem Zusammenhang dient das vorgeschaltete Real-Time IDS, (i) um Gegenmaßnahmen ohne hohe Verzögerungen durchzuführen (die Erzeugung von Flows und deren Export verzögert den Prozess der Angriffserkennung), (ii) da vor allem DoS / DDoS-Angriffe eine große Menge an Flow-Daten produzieren, welche die Flow-basierten IDS überlasten können und durch das Real-Time IDS reduziert werden können und (iii) im Regelfall insbesondere DDoS-Attacken hohe Auswirkungen auf den Betrieb des Backbone-Providers haben und deshalb effektiv verhindert werden müssen.

**Übrige IDSs:** Sollte der Manager eine weitergehende Analyse eines Angriffs initiieren, können die Protokoll-basierten/statistischen IDSs oder DPI-basierte IDSs aktiviert werden. Hierzu veranlasst der Manager eine Rekonfiguration des Routers, um nur den Teil des Datenverkehrs an die jeweiligen Systeme weiterzuleiten, der vorher bereits als verdächtig eingestuft wurde. Analog zu den Flow-basierten IDSs werden diese IDSs im Vorfeld auch über vorherige Alarmmeldungen informiert und melden ihre Analysen an den Manager. Wenn ein Angriff erkannt wurde, wird der Router durch den Manager angewiesen, den entsprechenden Datenverkehr zu blockieren. Sollte ein Angriff nicht bestätigt werden können, wird der Manager dafür sorgen, dass die verschiedenen IDSs (mit Ausnahme des Real-Time IDS sowie des/der Flow-basierte IDSs) den Datenverkehr für eine gewisse Zeit nicht erneut analysieren.

Zum genaueren Verständnis werden die einzelnen Schritte im Detail vorgestellt:

**Gewinnung von Indizien:** Sobald durch das Real-Time IDS bzw. das Flow-basierte IDS ein Indiz für einen möglichen Angriff entdeckt wurde, wird der entsprechende Datenstrom isoliert betrachtet. Da diese Untersuchung nicht auf der gesamten Nutzlast durchgeführt wird, wird die Privatsphäre der Nutzer respektiert und auf der anderen Seite ermöglicht dies die Verwendung preiswerter IDSs.

**Bewertung:** Basierend auf dem Alarm des Real-Time IDS, dem Flow-basierten IDS oder dem Protokollbasierten/statistischen IDS und dem entsprechenden CVE kann eine erste Einschätzung über das Ausmaß des Angriff unter Nutzung von CVSS durchgeführt werden; CVSS ist ein Industriestandard, der den Schweregrad der Schwachstellen in Computersystemen beschreibt. Ziel dieser Gewichtung ist das Herstellen einer Prioritätsreihenfolge für den Fall, dass aufgrund von Ressourcenlimitierungen nicht alle gemeldeten Angriffe im Detail untersucht werden können. Zusätzlich zu der CVSS-Gewichtung werden im Rahmen

der Architektur noch weitere (lokale) Kriterien verwendet, um die individuelle Bedrohung zu bewerten. Solche Kriterien sind beispielsweise, ob ein Angreifer in der Vergangenheit schon verdächtig war, oder ob sich das Ausmaß eines Angriffs (der derzeit untersucht wird) drastisch erhöht, bspw. bei Zunahme entsprechender Pakete, oder ob eine hohe Anzahl von ähnlichen Angriffen in der Vergangenheit beobachtet wurde. In diesem Zusammenhang wird auch Geolokalisation verwendet, was ebenfalls ein Bewertungskriterium zur Ähnlichkeit liefert. Dazu werden entsprechend Geo-Datenbanken eingesetzt. Je näher ein bereits erkanntes Indiz zu einem in der Vergangenheit erkannten Angriff ist, desto höher die Gewichtung. Geolokalisation hat hier auch noch einen zweiten Zweck: Abhängig von der Herkunft des Paketes kann ferner geprüft werden, ob evtl. ergriffenen Maßnahmen (z.B. die Anwendung von DPI) im Einklang mit den Gesetzen sind und so ein Provider nicht fälschlicherweise Gegenmaßnahmen ergreift, die für ihn juristische und vertragliche Konsequenzen nach sich ziehen.

**Detaillierte Analyse:** In Abhängigkeit des detektierten Angriffs und der vorangegangenen Wertung können nun spezifische IDSs zur erweiterten Analyse hinzugezogen werden. Dies ist von elementarer Bedeutung, da der Backbone-Provider die Glaubwürdigkeit und Vertraulichkeit gegenüber seinen Kunden riskiert und zugleich die rechtlichen und wirtschaftlichen Konsequenzen durch blockierte oder verworfene Datenströme nicht zu vernachlässigen sind. Aus diesem Grund möchte ein Backbone-Provider sehr sicher sein, bevor er eine Entscheidung wie das Blockieren des Datenstroms trifft. Dies kann auch bedeuten, dass ein Provider mehrere IDSs (z.B. mehrere verschiedene IDSs, wie bspw. Snort oder BRO) parallel für eine Untersuchung verwendet.

**Ergebnisbewertung:** Nachdem alle Untersuchungen beendet sind, erfolgt die Korrelation und Bewertung der Ergebnisse, um ein möglichst hohes Maß an Verlässlichkeit gewährleisten zu können, bevor mögliche Gegenmaßnahmen ergriffen werden. Dies ist insbesondere dann relevant, wenn mehrere IDSs zur Untersuchung verwendet wurden.

**Reaktion und Gegenmaßnahmen:** Der letzte Schritt ist die Einleitung einer Reaktion auf die erfolgte Detektion. Dies kann durch entsprechende Gegenmaßnahmen, wie das Verwerfen von Paketen und/oder der Rekonfiguration des Core Routers, von staten gehen. Zugleich ist eine Weitermeldung an andere Router (beziehungsweise Instanzen) und die lokale Speicherung des Vorfalls möglich.

## 4 Prototypische Umsetzung

Zur Evaluation der Leistungsfähigkeit der Architektur wird diese derzeit prototypisch umgesetzt. Der Großteil der Managementkomponente wird aufgrund der Performance in C realisiert; hierbei kommen Open Source Bibliotheken zum Einsatz, um die Entwicklungszeit zu verkürzen. Als Datenbankbackend kommt MySQL zum Einsatz. Durch diverse APIs können Daten wie bspw. CVE importiert werden. In diesem ersten Prototyp werden

als IDSs im Wesentlichen eine speziell für Protokollanalyse angepasste, als auch eine in der Standardkonfiguration befindliche Version von Snort verwendet. Der bereits in erster Version verfügbare Proof of Concept wird derzeit in das Labor des Institutes integriert, die ersten intensiven Tests laufen bereits. Das Institutsnetz besteht hierbei aus den sicherheitstechnisch strikt voneinander getrennten Produktiv-, Forschungs- sowie Büronetzen. Als Geolokalisationslösung kommt hier eine von uns speziell entwickelte Software zum Einsatz, welche IP Adressen auf Landesebene mit einer Genauigkeit von 99.78% erkennt (vgl. [KGR13a, KGR13b]).

Um die vorgestellte Architektur mit dem derzeitigen „State of the Art“ zu vergleichen, werden verschiedene Kriterien, wie die Genauigkeit und die Rate von False Positives, herangezogen. Dabei ist es nicht das Ziel möglichst viele Vorfälle zu detektieren, sondern die False Positives zu reduzieren.

Hierzu wird parallel zu unserer Architektur eine vergleichende Betrachtung mit „State of the Art“-Systemen durchgeführt. Die Anbindungen der einzelnen Komponenten und die zugrundeliegenden Bedingungen, wie zum Beispiel beim Export der Flows, sind in beiden Umgebungen gleich (z.B. zwei identische Cisco Catalyst 6509-E Router).

## **5 Zusammenfassung und Ausblick**

Die globale Vernetzung und die Durchdringung des alltäglichen Lebens durch IKT führen zu einer wachsenden Zahl sicherheitstechnischer Herausforderungen, derer sich auch Backbone-Provider nicht verschließen können. Im Rahmen der vorliegenden Publikation wurde ein mehrstufiger Ansatz zur Angriffserkennung in Hochgeschwindigkeits-Backbone-Netzen vorgestellt. Dieser ermöglicht es, neben den finanziellen und technischen auch rechtliche Aspekte, beispielsweise in Bezug auf den Datenschutz sowie die Privatsphäre, zu berücksichtigen. Obwohl sich die Fähigkeit des Gesamtsystems auf die der Flow-basierenden Verfahren beschränkt, ist die Architektur gerade deswegen für Backbone-Betreiber geeignet. Denn diese Verfahren ermöglichen eine vergleichsweise kostengünstige Integration in bestehende Umgebungen und effektive Behandlung von Angriffen, die von besonderem Interesse für Backbone-Betreiber sind.

Zukünftige Arbeiten werden sich mit der Anpassung und Erweiterung der Architektur befassen, um so eine möglichst weite Verbreitung in unterschiedlichen Netzen und einer automatisierten Anpassung an diese zu ermöglichen. Hierfür sind allerdings mehr Informationen zu den technischen und rechtlichen Gegebenheiten unterschiedlicher Umgebungen von Nöten. Weitere Gedanken befassen sich mit einer noch weitergehenden Integration von Geolokalisationsmechanismen in die vorgestellte Architektur, sowie eine tiefergehenden Betrachtung (*Evaluierung*) der Abhängigkeit zwischen Detektion und Geolokalisation. Dies kann vor allem der länderspezifischen Behandlung von Vorfällen, wie auch der vorgestellten Bewertung derer zuträglich sein.



## Danksagung

Diese Arbeit wurde teilweise durch Flamingo, einem Network of Excellence Projekt (ICT-318488) des 7. EU-Forschungsrahmenprogramms der europäischen Kommission, gefördert.

## Literatur

- [Arb] Arbor Networks. Worldwide Infrastructure Security Report - 2012 Volume VIII. [http://pages.arbornetworks.com/rs/arbor/images/WISR2012\\_EN.pdf](http://pages.arbornetworks.com/rs/arbor/images/WISR2012_EN.pdf), zuletzt besucht am 14.04.2014.
- [CTA13] Benoit Claise, Brian Trammell und Paul Aitken. Specification of the IP Flow Information Export (IPFIX) Protocol for the Exchange of Flow Information. RFC 7011 (Internet Standard), 2013.
- [ES10] Patricia Takako Endo und Djamel Fawzi Hadj Sadok. Whois Based Geolocation: A Strategy to Geolocate Internet Hosts. In *Proceedings of the 2010 24th IEEE International Conference on Advanced Information Networking and Applications*, Seiten 408–413, 2010.
- [HBSP13] Rick Hofstede, Vaclav Bartos, Anna Sperotto und Aiko Pras. Towards Real-Time Intrusion Detection for NetFlow and IPFIX. In *Proceedings of the 9th International Conference on Network and Service Management, CNSM'13*, Seiten 227–234, 2013.
- [HFc11] B. Huffaker, M. Fomenkov und k. claffy. Geocompare: a comparison of public and commercial geolocation databases - Technical Report. Bericht, CAIDA, May 2011.
- [KGR13a] Robert Koch, Mario Golling und Gabi Dreo Rodosek. Advanced Geolocation of IP Addresses. In *International Conference on Communication and Network Security (ICCNS)*, Seiten 1–10, 2013.
- [KGR13b] Robert Koch, Mario Golling und Gabi Dreo Rodosek. Geolocation and Verification of IP-Addresses with Specific Focus on IPv6. In *5th International Symposium on Cyberspace Safety and Security (CSS 2013)*, Seiten 1–20. Springer, 2013.
- [Man13] Mandiant. APT1 - Exposing One of China's Cyber Espionage Units, 2013. <http://intelreport.mandiant.com/>, zuletzt besucht am 14.04.2014.
- [SM10] Karen Scarfone und Peter Mell. Intrusion detection and prevention systems. In *Handbook of Information and Communication Security*, Seiten 177–192. Springer, 2010.
- [SSS<sup>+</sup>10] Anna Sperotto, Gregor Schaffrath, Ramin Sadre, Cristian Morariu, Aiko Pras und Burkhard Stiller. An Overview of IP Flow-Based Intrusion Detection. *IEEE Communications Surveys & Tutorials*, 12(3):343–356, 2010.
- [ZFdRD05] Artur Ziviani, Serge Fdida, José F. de Rezende und Otto Carlos M. B. Duarte. Improving the Accuracy of Measurement-based Geographic Location of Internet Hosts. *Comput. Netw. ISDN Syst.*, 47(4):503–523, Marz 2005.

## *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 sensor-gestützte Chirurgie; Workshop des SFB 414.
- P-5 Andy Schürr (Hg.): OMER – Object-Oriented Modeling of Embedded Real-Time Systems.
- P-6 Hans-Jürgen Appelpoth, 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 Arbeits-tagung.
- 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 Verletz-barkeit 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 (Ergänzungs-band).
- 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 – Middle-ware 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 Fähnrich, 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 1<sup>st</sup> 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 3<sup>rd</sup> 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 1<sup>st</sup> International Working Conference on Business Process and Services Computing BPSC 2007
- P-117 Ryszard Kowalczyk (Ed.): Grid service engineering and management The 4<sup>th</sup> 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 Hermann, 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.)  
9<sup>th</sup> 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.)  
10<sup>th</sup> 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.)  
3<sup>rd</sup> International Conference on Electronic Voting 2008  
Co-organized by Council of Europe, Gesellschaft für Informatik und 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.)  
9<sup>th</sup> International Conference on Innovative Internet Community Systems  
I<sup>2</sup>CS 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 Kirm, 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 3<sup>rd</sup> 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. Claudepein, 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.)  
10<sup>th</sup> International Conference on Innovative Internet Community Systems (I<sup>2</sup>CS) – 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.)  
4<sup>th</sup> 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 Juhrisch (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 Informationssysteme 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.)  
6<sup>th</sup> 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-186 Gerald Eichler, Axel Küpper, Volkmar Schau, Hacène Fouchal, Herwig Unger (Eds.)  
11<sup>th</sup> International Conference on Innovative Internet Community Systems (I<sup>2</sup>CS)
- P-187 Paul Müller, Bernhard Neumair, Gabi Dreö Rodosek (Hrsg.)  
4. DFN-Forum Kommunikationstechnologien, Beiträge der Fachtagung 20. Juni bis 21. Juni 2011 Bonn
- P-188 Holger Rohland, Andrea Kienle, Steffen Friedrich (Hrsg.)  
DeLFI 2011 – Die 9. e-Learning Fachtagung Informatik der Gesellschaft für Informatik e.V. 5.–8. September 2011, Dresden
- P-189 Thomas, Marco (Hrsg.)  
Informatik in Bildung und Beruf INFOS 2011  
14. GI-Fachtagung Informatik und Schule
- P-190 Markus Nüttgens, Oliver Thomas, Barbara Weber (Eds.)  
Enterprise Modelling and Information Systems Architectures (EMISA 2011)
- P-191 Arslan Brömme, Christoph Busch (Eds.)  
BIOSIG 2011  
International Conference of the Biometrics Special Interest Group
- P-192 Hans-Ulrich Heiß, Peter Pepper, Holger Schlingloff, Jörg Schneider (Hrsg.)  
INFORMATIK 2011  
Informatik schafft Communities
- P-193 Wolfgang Lehner, Gunther Piller (Hrsg.)  
IMDM 2011
- P-194 M. Clasen, G. Fröhlich, H. Bernhardt, K. Hildebrand, B. Theuvsen (Hrsg.)  
Informationstechnologie für eine nachhaltige Landbewirtschaftung Fokus Forstwirtschaft
- P-195 Neeraj Suri, Michael Waidner (Hrsg.)  
Sicherheit 2012  
Sicherheit, Schutz und Zuverlässigkeit  
Beiträge der 6. Jahrestagung des Fachbereichs Sicherheit der Gesellschaft für Informatik e.V. (GI)
- P-196 Arslan Brömme, Christoph Busch (Eds.)  
BIOSIG 2012  
Proceedings of the 11<sup>th</sup> International Conference of the Biometrics Special Interest Group
- P-197 Jörn von Lucke, Christian P. Geiger, Siegfried Kaiser, Erich Schweighofer, Maria A. Wimmer (Hrsg.)  
Auf dem Weg zu einer offenen, smarten und vernetzten Verwaltungskultur  
Gemeinsame Fachtagung Verwaltungsinformatik (FTVI) und Fachtagung Rechtsinformatik (FTRI) 2012
- P-198 Stefan Jähnichen, Axel Küpper, Sahin Albayrak (Hrsg.)  
Software Engineering 2012  
Fachtagung des GI-Fachbereichs Softwaretechnik
- P-199 Stefan Jähnichen, Bernhard Rumpe, Holger Schlingloff (Hrsg.)  
Software Engineering 2012  
Workshopband
- P-200 Gero Mühl, Jan Richling, Andreas Herkersdorf (Hrsg.)  
ARCS 2012 Workshops
- P-201 Elmar J. Sinz Andy Schürr (Hrsg.)  
Modellierung 2012
- P-202 Andrea Back, Markus Bick, Martin Breunig, Key Poustchi, Frédéric Thiesse (Hrsg.)  
MMS 2012: Mobile und Ubiquitäre Informationssysteme
- P-203 Paul Müller, Bernhard Neumair, Helmut Reiser, Gabi Dreö Rodosek (Hrsg.)  
5. DFN-Forum Kommunikationstechnologien  
Beiträge der Fachtagung
- P-204 Gerald Eichler, Leendert W. M. Wienhofen, Anders Kofod-Petersen, Herwig Unger (Eds.)  
12<sup>th</sup> International Conference on Innovative Internet Community Systems (I<sup>2</sup>CS 2012)
- P-205 Manuel J. Kripp, Melanie Volkamer, Rüdiger Grimm (Eds.)  
5<sup>th</sup> International Conference on Electronic Voting 2012 (EVOTE2012)  
Co-organized by the Council of Europe, Gesellschaft für Informatik und E-Voting.CC
- P-206 Stefanie Rinderle-Ma, Mathias Weske (Hrsg.)  
EMISA 2012  
Der Mensch im Zentrum der Modellierung
- P-207 Jörg Desel, Jörg M. Haake, Christian Spannagel (Hrsg.)  
DeLFI 2012: Die 10. e-Learning Fachtagung Informatik der Gesellschaft für Informatik e.V.  
24.–26. September 2012

- P-208 Ursula Goltz, Marcus Magnor, Hans-Jürgen Appelrath, Herbert Matthies, Wolf-Tilo Balke, Lars Wolf (Hrsg.)  
INFORMATIK 2012
- P-209 Hans Brandt-Pook, André Fleer, Thorsten Spitta, Malte Wattenberg (Hrsg.)  
Nachhaltiges Software Management
- P-210 Erhard Plödereder, Peter Dencker, Herbert Klenk, Hubert B. Keller, Silke Spitzer (Hrsg.)  
Automotive – Safety & Security 2012  
Sicherheit und Zuverlässigkeit für automobile Informationstechnik
- P-211 M. Clasen, K. C. Kersebaum, A. Meyer-Aurich, B. Theuvsen (Hrsg.)  
Massendatenmanagement in der Agrar- und Ernährungswirtschaft  
Erhebung - Verarbeitung - Nutzung  
Referate der 33. GIL-Jahrestagung  
20. – 21. Februar 2013, Potsdam
- P-212 Arslan Brömme, Christoph Busch (Eds.)  
BIOSIG 2013  
Proceedings of the 12th International Conference of the Biometrics Special Interest Group  
04.–06. September 2013  
Darmstadt, Germany
- P-213 Stefan Kowalewski, Bernhard Rumpe (Hrsg.)  
Software Engineering 2013  
Fachtagung des GI-Fachbereichs Softwaretechnik
- P-214 Volker Markl, Gunter Saake, Kai-Uwe Sattler, Gregor Hackenbroich, Bernhard Mitschang, Theo Härder, Veit Köppen (Hrsg.)  
Datenbanksysteme für Business, Technologie und Web (BTW) 2013  
13. – 15. März 2013, Magdeburg
- P-215 Stefan Wagner, Horst Lichter (Hrsg.)  
Software Engineering 2013  
Workshopband  
(inkl. Doktorandensymposium)  
26. Februar – 1. März 2013, Aachen
- P-216 Gunter Saake, Andreas Henrich, Wolfgang Lehner, Thomas Neumann, Veit Köppen (Hrsg.)  
Datenbanksysteme für Business, Technologie und Web (BTW) 2013 – Workshopband  
11. – 12. März 2013, Magdeburg
- P-217 Paul Müller, Bernhard Neumair, Helmut Reiser, Gabi Dreö Rodosek (Hrsg.)  
6. DFN-Forum Kommunikationstechnologien  
Beiträge der Fachtagung  
03.–04. Juni 2013, Erlangen
- P-218 Andreas Breiter, Christoph Rensing (Hrsg.)  
DeLFI 2013: Die 11 e-Learning Fachtagung Informatik der Gesellschaft für Informatik e.V. (GI)  
8. – 11. September 2013, Bremen
- P-219 Norbert Breier, Peer Stechert, Thomas Wilke (Hrsg.)  
Informatik erweitert Horizonte  
INFOS 2013  
15. GI-Fachtagung Informatik und Schule  
26. – 28. September 2013
- P-220 Matthias Horbach (Hrsg.)  
INFORMATIK 2013  
Informatik angepasst an Mensch, Organisation und Umwelt  
16. – 20. September 2013, Koblenz
- P-221 Maria A. Wimmer, Marijn Janssen, Ann Macintosh, Hans Jochen Scholl, Efthimios Tambouris (Eds.)  
Electronic Government and Electronic Participation  
Joint Proceedings of Ongoing Research of IFIP EGOV and IFIP ePart 2013  
16. – 19. September 2013, Koblenz
- P-222 Reinhard Jung, Manfred Reichert (Eds.)  
Enterprise Modelling and Information Systems Architectures (EMISA 2013)  
St. Gallen, Switzerland  
September 5. – 6. 2013
- P-223 Detlef Hühnlein, Heiko Roßnagel (Hrsg.)  
Open Identity Summit 2013  
10. – 11. September 2013  
Kloster Banz, Germany
- P-224 Eckhart Hanser, Martin Mikusz, Masud Fazal-Baqaie (Hrsg.)  
Vorgehensmodelle 2013  
Vorgehensmodelle – Anspruch und Wirklichkeit  
20. Tagung der Fachgruppe Vorgehensmodelle im Fachgebiet Wirtschaftsinformatik (WI-VM) der Gesellschaft für Informatik e.V.  
Lörrach, 2013
- P-225 Hans-Georg Fill, Dimitris Karagiannis, Ulrich Reimer (Hrsg.)  
Modellierung 2014  
19. – 21. März 2014, Wien
- P-226 M. Clasen, M. Hamer, S. Lehnert, B. Petersen, B. Theuvsen (Hrsg.)  
IT-Standards in der Agrar- und Ernährungswirtschaft Fokus: Risiko- und Krisenmanagement  
Referate der 34. GIL-Jahrestagung  
24. – 25. Februar 2014, Bonn

- P-227 Wilhelm Hasselbring,  
Nils Christian Ehmke (Hrsg.)  
Software Engineering 2014  
Fachtagung des GI-Fachbereichs  
Softwaretechnik  
25. – 28. Februar 2014  
Kiel, Deutschland
- P-228 Stefan Katzenbeisser, Volkmar Lotz,  
Edgar Weippl (Hrsg.)  
Sicherheit 2014  
Sicherheit, Schutz und Zuverlässigkeit  
Beiträge der 7. Jahrestagung des  
Fachbereichs Sicherheit der  
Gesellschaft für Informatik e.V. (GI)  
19.–21. März 2014, Wien
- P-231 Paul Müller, Bernhard Neumair,  
Helmut Reiser, Gabi Dreö Rodosek  
(Hrsg.)  
7. DFN-Forum  
Kommunikationstechnologien  
16.–17. Juni 2014  
Fulda

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](mailto:druckverlag@koellen.de)



